

John D. Barrow

100 COISAS ESSENCIAIS QUE NÃO SABIA QUE NÃO SABIA

Perceba o Mundo Através da Matemática



VISIT...

LANZAROTE
Caliente.COM

Título:

100 Coisas essenciais que não sabia que não sabia

Título original:

100 Essential Things You Didn't Know You Didn't Know

Autor:

John D. Barrow

Tradução:

Rita Figueiredo

Revisão:

Ana Pedro

Capa:

Estúdios Horizonte

Imagens da capa:

© Petr Václavěk/Forolia



© Livros Horizonte, 2011

© John D. Barrow, 2008

Primeira edição publicada no Reino Unido por: The Bodley Head

ISBN 978-972-24-1754-9

Paginação:

Estúdios Horizonte

Impressão:

Guide

Abril de 2013

Dep. Legal n.º 358281/13



Esta edição segue a grafia do Novo Acordo Ortográfico da Língua Portuguesa.

Reservados todos os direitos de publicação

total ou parcial para a língua portuguesa por

LIVROS HORIZONTE, LDA.

Rua das Chagas, 17-1.º Dt.º – 1200-106 LISBOA

E-mail: geral@livroshorizonte.pt

www.livroshorizonte.pt

ÍNDICE

Prefácio	13
1 Poste do mês	15
2 Uma noção de equilíbrio	18
3 Macacadas	20
4 Dia da Independência	23
5 Râguebi e relatividade	25
6 Numa roda-viva	27
7 Uma noção de proporção	29
8 Por que é que a outra fila avança sempre mais depressa?	32
9 Dois é bom, três é de mais	34
10 O mundo é mesmo pequeno	37
11 Estabelecer pontes	39
12 Nas cartas	42
13 Da numeração romana ao críquete	46
14 Relações	49
15 As certezas das corridas	52
16 Saltos em altura	54
17 Superficialidade	57
18 IVA até à eternidade	60
19 Viver numa simulação	62
20 Emergência	66
21 Como empurrar um carro	68
22 Feedback positivo	70
23 O passo do bêbedo	72
24 O fingimento	74

25 A falha das médias	77
26 O origami do universo	79
27 Problemas fáceis e difíceis	81
28 Isto é um recorde?	84
29 "Lotaria faça-você-mesmo"	87
30 Não acredito!	89
31 Fogo!	91
32 O problema da secretária	93
33 Divórcio amigável: a solução vantajosa para ambas as partes	96
34 Feliz aniversário	99
35 A inclinação dos moinhos de vento	102
36 Prestidigitação verbal	105
37 Investimentos financeiros com viajantes no tempo	107
38 Um pensamento pelos seus tostões	110
39 Decomposição da lei das médias	113
40 Quanto tempo é provável que as coisas sobrevivam?	115
41 O presidente que preferia o triângulo ao pentágono	118
42 Códigos secretos no seu bolso	121
43 Tenho uma péssima memória para nomes	125
44 O cálculo prolonga a vida	127
45 Bater asas	129
46 O seu número foi escolhido	132
47 Duplique o seu dinheiro	134
48 Reflexos de rostos	135
49 O matemático mais infame	138
50 Montanhas-russas e saídas das autoestradas	141
51 Uma explosão por medida	145
52 Caminhe, não corra	148
53 Ler pensamentos	150
54 O planeta dos enganadores	152
55 Como ganhar a lotaria	154
56 Um jogo de futebol bizarro	157
57 Um problema de arcos	159
58 Contar por oitos	161
59 Uma maioria representativa	163

60 O campeonato de duas cabeças	165
61 Criar algo a partir do nada	167
62 Como manipular uma eleição	169
63 O balanço do pêndulo	172
64 Uma bicicleta com rodas quadradas	174
65 De quantos guardas precisa uma galeria de arte?	176
66 ... Então é uma prisão?	180
67 Um truque de <i>snooker</i>	182
68 Irmãos e irmãs	185
69 Jogar limpo com uma moeda viciada	188
70 As maravilhas da tautologia	190
71 Raquetadas	192
72 Fazer as malas	195
73 Fazer novamente as malas	197
74 Tigre agachado	202
75 As manchas do leopardo	204
76 A loucura das multidões	206
77 O homem dos diamantes	209
78 As três leis da robótica	213
79 Romper com as convenções	216
80 Google nas caraíbas: o poder da matriz	219
81 Aversão à perda	223
82 Na ponta do lápis	225
83 Teste de resistência do esparguete	227
84 O Gherkin	229
85 Calcular o índice de preços	232
86 A onisciência pode ser uma desvantagem	235
87 Por que é que as pessoas não são mais espertas	237
88 O homem do metro	238
89 Não há números desinteressantes	240
90 Incógnito	242
91 O paradoxo da patinagem no gelo	244
92 A regra de dois	247
93 Segregação e "micromotivos"	249
94 Não seguir a corrente	251
95 A lógica de Venn	253

96 Alguns benefícios da irracionalidade	255
97 Fórmulas estranhas	260
98 Caos	263
99 Todos a bordo	265
100 A aldeia global	268
Notas	269

Continuei a estudar aritmética com o meu pai, passando orgulhosamente das frações aos decimais. Acabei por chegar a um ponto em que um determinado número de vacas comia uma certa quantidade de erva e em que tanques se enchiam de água num determinado número de horas. Parecia-me absolutamente cativante.

Agatha Christie

Para David e Emma

PREFÁCIO

Este é um livro composto por fragmentos dispersos – pequenas histórias acerca de aplicações invulgares da matemática à vida quotidiana e de outras coisas não muito diferentes. São cem à escolha, não obedecem a nenhuma ordem em particular, não têm interesses ocultos nem um fio condutor invisível. Por vezes o leitor encontrará apenas palavras, mas haverá outras em que também encontrará números e, muito ocasionalmente, algumas notas que lhe apresentarão as fórmulas por trás das aparências. A matemática é interessante e importante porque pode ensinar-lhe coisas acerca do mundo que não podem ser aprendidas por outro meio. No que diz respeito às profundezas da física fundamental ou à dimensão do universo astronómico, já nos habituámos a esperar a presença da matemática. No entanto, espero que com este livro descubra como ideias simples podem dar-lhe uma nova perspetiva de coisas que de outra forma lhe pareceriam aborrecidamente familiares ou que passariam simplesmente despercebidas.

Muitos dos exemplos contidos nas páginas que se seguem foram inspirados pelo Millenium Mathematics Project*, que dirijo em Cambridge desde 1999. Quando cumprido, o desafio de mostrar que a matemática tem algo a ensinar-nos acerca de tudo no mundo que nos rodeia pode desempenhar um papel importante na motivação e informação das pessoas, jovens e velhas, e ensiná-las a apreciar e a entender o papel fundamental da matemática na nossa compreensão do mundo.

Gostaria de agradecer a Steve Brams, Marianne Freiberger, Jenny Gage, John Haigh, Jörg Hensgen, Helen Joyce, Tom Körner, Imre Leader,

* www.mmp.maths.org

Drummond Moir, Robert Osserman, Jenny Piggott, David Spiegelhalter, Will Sulkin, Rachel Thomas, John W. Webb, Marc West e Robin Wilson pelas úteis trocas de ideias, pelo seu encorajamento e por outras contribuições práticas para a compilação final de coisas essenciais que o leitor tem agora à sua frente.

Finalmente, quero agradecer a Elizabeth, David, Roger e Louise pelo seu maravilhoso e desconcertante interesse por este livro. Muitos destes meus familiares explicam-me agora por que motivo os postes elétricos são compostos por triângulos e por que é que os equilibristas transportam uma vara comprida nas mãos. Em breve, também o leitor saberá a resposta.

John D. Barrow

Agosto de 2008, Cambridge

1 | POSTE DO MÊS

Tal como Moisés a dividir as águas, o cabo de alta tensão 4YG8 da National Grid Company guia os postes seus companheiros através desta zona residencial de Oxfordshire em direção à "terra prometida" da Central Elétrica de Didcot.

Pylon of the Month de dezembro de 1999

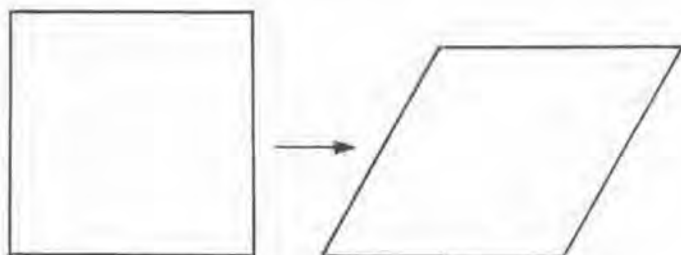
Há alguns *websites* fascinantes por aí, mas nenhum é tão encantador quanto o emblemático *Pylon of the Month*^{*}, que em tempos se dedicou a divulgar imagens dos postes elétricos mais belos e sedutores do mundo. Os que poderá ver no *website* indicado abaixo ficam na Escócia. Infelizmente, o *Pylon of the Month* parece ter-se tornado um *website* fantasma, mas ainda podemos aprender alguma coisa com ele, já que, para um matemático, cada poste conta uma história. São tão imponentes e a sua presença é de tal forma ubíqua que, à semelhança da gravidade, quase passam despercebidos.

Da próxima vez que viajar de comboio, observe atentamente os postes elétricos enquanto passam a grande velocidade pelas janelas do comboio. Cada um deles é composto por uma rede de barras metálicas que usa uma forma poligonal única e recorrente. Essa forma é o triângulo. Dentro deles estão encaixados triângulos grandes e outros mais pequenos. Até mesmo os que parecem quadrados e retângulos são apenas pares de triângulos. O motivo pelo qual têm esta configuração faz parte de uma interessante história matemática que começou no início do século XIX, com o trabalho do matemático francês Augustin-Louis Cauchy.

De todas as formas poligonais que poderíamos criar unindo barras metálicas, o triângulo é especial. É a única forma que é *rígida*. Se colocar-

^{*} <http://www.drookitagain.co.uk/coppermine/thumbnails.php?album=34>

mos articulações nos cantos de outro polígono, podemos gradualmente mudar a sua forma, originando uma nova sem que o metal se curve. Uma forma quadrada ou retangular proporciona-nos um exemplo simples: vemos que pode ser deformada de modo a criar um paralelogramo sem dobrar o metal. Esta consideração é importante se quisermos manter a estabilidade estrutural face aos ventos e às mudanças de temperatura. É por isso que os postes elétricos parecem sempre grandes totens ao deus de todos os triângulos.



Se avançarmos para as formas tridimensionais, a situação é bastante diferente: Cauchy mostrou que *todos* os poliedros convexos (ou seja, aqueles que têm todas as faces voltadas para fora) com faces rígidas e com os cantos articulados são rígidos. E, de facto, o mesmo é verdade em relação aos poliedros convexos em espaços com quatro ou mais dimensões.

Então e os poliedros não-convexos, que podem ter algumas faces apontadas para dentro? Parecem ter mais probabilidades de se abaterem sobre si mesmos. Esta questão permaneceu em aberto até 1978, quando Robert Connelly encontrou um exemplo com faces não-convexas e que não era rígido, e depois demonstrou que em todos os casos as mudanças de flexibilidade possíveis mantinham o volume total do poliedro igual. Contudo, os exemplos de poliedros não-convexos existentes, ou que possam vir a ser encontrados no futuro, parecem não ter qualquer interesse prático imediato para os engenheiros estruturais, uma vez que são especiais na medida em que exigem uma construção perfeitamente exata, como equilibrar uma agulha sobre a sua ponta. Qualquer desvio a este cenário cria um exemplo rígido. É por isso que os matemáticos dizem que “quase todos” os poliedros são rígidos. Tudo isto

leva a crer que a estabilidade estrutural é fácil de alcançar. Ainda assim, os postes elétricos cedem e caem. De certeza que consegue perceber porquê.



2 | UMA NOÇÃO DE EQUILÍBRIO

Apesar da minha educação privilegiada, sou muito equilibrado. Tenho partes iguais de ressentimento e de sentimento de inferioridade.

Russell Crowe, in *Uma Mente Brilhante*

Independentemente do que faça na vida, haverá alturas em que sentirá que está a caminhar na corda bamba entre o sucesso e o fracasso, a tentar equilibrar várias coisas ou evitar que uma determinada atividade ocupe todos os seus minutos de tempo livre. Mas e no caso das pessoas que caminham realmente na corda bamba? No outro dia, vi num filme antigo uma imagem que nos é muito familiar: um funâmbulo louco desafiava a morte, atravessando uma corda bamba sobre uma ravina e um rio turbulento. Bastava-lhe dar um passo em falso para se tornar mais uma vítima da lei da gravidade de Newton.

Todos nós já tentámos equilibrar-nos em cima de degraus ou de tábuas de madeira, e a experiência ensinou-nos que há algumas coisas que podem ajudar-nos a manter-nos direitos e equilibrados: não nos inclinarmos para fora do centro, mantermos o corpo direito e o nosso centro de gravidade baixo. Todas essas coisas que se ensinam na escola de artes circenses. Mas aqueles funâmbulos parecem levar sempre uma vara muito comprida nas mãos. Às vezes a vara até se curva por causa do peso, e outras vezes chega mesmo a ter pesados baldes presos a cada uma das extremidades. Por que acha que os funâmbulos fazem isso?

O conceito-chave que é necessário compreender e que explica a utilização dessa vara comprida por parte dos funâmbulos é a inércia. Quanto maior for a inércia, mas lentamente se deslocará quando uma força for aplicada. Não tem nada que ver com o centro de gravidade. Quanto mais longe do centro estiver distribuída a massa, mais elevada será a inércia

do corpo e mais difícil será movê-lo. Tomemos como exemplo duas esferas de materiais diferentes que têm o mesmo diâmetro e a mesma massa, uma sólida e uma oca. Verificamos que é a esfera oca, com a massa mais distante, distribuída pela sua superfície, que desce mais lentamente uma superfície inclinada. Da mesma forma, transportar uma vara comprida aumenta a inércia do funambulo, afastando a massa da linha central do corpo — a inércia é igual à massa multiplicada pela distância ao quadrado. Consequentemente, quaisquer pequenas oscilações da posição de equilíbrio acontecem mais devagar. O período de oscilação é mais longo e o funambulo tem mais tempo para reagir às oscilações e para recuperar o equilíbrio. Faça a experiência e veja como é mais fácil equilibrar uma vara de um metro na ponta de um dedo, do que uma com apenas dez centímetros.

3 | MACACADAS

Tenho um corretor horto gráfico
Que veio com o meu computador
Marca clara mente para revi são
Os erros que não com sigo ver

Eu sei-o neste poema
E estou certo que gostaram de saber
Que não tem um único error
Foio que disse o meu corretor...

Barri Haynes

A imagem lendária de um exército de macacos a digitar letras ao acaso e a acabar por produzir as obras completas de Shakespeare parece ter emergido gradualmente ao longo de muito tempo. N' *As Viagens de Gulliver*, livro escrito em 1726, Jonathan Swift fala-nos do mítico professor da Grande Academia de Lagado que tenta criar um catálogo completo de todo o conhecimento científico fazendo os seus alunos produzirem continuamente sequências aleatórias de letras por meio de um dispositivo mecânico de impressão. A primeira máquina de escrever mecânica tinha sido patenteada em 1714. Depois de vários matemáticos franceses dos séculos XVIII e XIX terem usado o exemplo de uma grande obra composta por uma torrente aleatória de letras originadas por um dispositivo de impressão como ilustrativo de uma improbabilidade extrema, surgiu o exemplo dos macacos, em 1909, quando o matemático francês Émile Borel sugeriu que um grupo de macacos a digitar letras aleatoriamente acabaria por produzir todos os livros contidos na *Bibliothèque Nationale*. Arthur Eddington usou esta analogia no seu famoso livro *The Nature of the Physical World*, em 1928, anglicizando a biblioteca em questão: "Se eu

deixar que os meus dedos divaguem indolentemente pelas teclas de uma máquina de escrever, *pode* dar-se o caso de a sequência por mim originada constituir uma frase inteligível. Se um exército de macacos dedilhasse aleatoriamente várias máquinas de escrever, *poderia* acabar por escrever todos os livros contidos no Museu Britânico.”

Este exemplo tão repetido ao longo dos tempos acabou por optar pelas *Obras Completas de Shakespeare* como principal candidato à recriação aleatória. Curiosamente, houve um *website* que simulou a escrita aleatória e contínua numa máquina de escrever e depois confrontou os resultados com as *Obras Completas de Shakespeare* para identificar as correspondências. A simulação das ações dos macacos começou no dia 1 de julho de 2003 com cem macacos e a população de macacos foi sendo duplicada de poucos em poucos dias até recentemente. Durante esse tempo, produziram mais de 10^{35} páginas, cada uma das quais com 2000 caracteres.

Foi mantido um registo contínuo das sequências diárias e globais até que o projeto de simulação dos macacos de Shakespeare parou de fazer atualizações em 2007. Os registos diários são razoavelmente estáveis, com cada sequência a rondar os 18 ou 19 caracteres, e registando uma tendência constante para aumentar. Por exemplo, uma das sequências de 18 caracteres que os macacos geraram está contida no seguinte fragmento:

... Theseus. Now faire UWfllaNWSK2d6L;wb...

Os primeiros 18 caracteres correspondem a parte de um excerto de *Sonho de uma noite de verão* onde se lê:

... us. Now faire Hippolita, our nuptiall houre...

Durante algum tempo, as sequências obtidas tiveram 21 caracteres de comprimento, originando:

... KING. Let fame, that wtIA"yh!"VYONovwsFOsbhzkLH...

que tem correspondência com 21 caracteres da peça *Penas de amor perdidas*:

KING. Let fame, that all hunt after in their lives,
Live regist'ed upon our brazen tombs,
And then grace us in the disgrace of death; ...

Em dezembro de 2004, o registo alcançou uma correspondência de 23 caracteres com:

Poet. Good day Sir FhlOiX5a]OM.MLGtUGSxX4lfeHQbktQ...

que tinha correspondência com parte da peça *Timão de Atenas*:

Poet. Good day Sir

Pain. I am glad y'are well

Poet. I have not seene you long, how goes the World?

Pain. It weares sir, as it growes...

Em janeiro de 2005, após o equivalente a 2 737 850 milhões de milhões de milhões de milhões de milhões de milhões de anos de escrita aleatória por parte dos macacos, o registo alcançou uma correspondência de 24 caracteres, com:

RUMOR. Open your ears; 9r"5j5&?OWTY Z0d'B-nEoF.vjSqj[...

que corresponde a 24 letras de *Henrique IV Parte 2*:

RUMOUR. Open your ears; for which of you will stop

The vent of hearing when loud Rumour speaks?

Como vê, é tudo uma questão de tempo!

4 | DIA DA INDEPENDÊNCIA

Li que existe uma probabilidade de 1 em 1000 de uma pessoa entrar num avião em que existe uma bomba. Assim sendo, comecei a levar uma bomba sempre que andava de avião; calculei que as probabilidades de haver duas pessoas com bombas num avião seriam astronómicas.

Anónimo

A comemoração do Dia da Independência (4 de julho) de 1977 é uma data da qual me lembro bem. Para além de ter sido um dos dias mais quentes que houve em Inglaterra em muitos anos, foi o dia da defesa da minha tese em Oxford. A independência, ainda que de um tipo ligeiramente diferente, revelou-se extremamente importante, porque a primeira pergunta que os examinadores me fizeram não tinha nada que ver com cosmologia, que era o tema da minha tese. Em vez disso, foi acerca de estatística. Um dos examinadores tinha encontrado 32 erros tipográficos na tese (naquela altura não havia processadores de texto e aos corretores ortográficos). O outro examinador tinha encontrado 23. A pergunta era: quantos mais haveria que não tinham sido encontrados por nenhum deles? Ao fim de algum tempo a analisar várias folhas de papel, concluiu-se que 16 dos erros tinham sido encontrados por ambos os examinadores. É surpreendente que esta informação nos permita responder à pergunta desde que se parta da suposição de que os dois examinadores trabalham separadamente, de modo a que a probabilidade de um deles encontrar um erro não seja afetada pelo facto de o outro encontrar um erro ou não.

Suponhamos que os dois examinadores encontraram A e B erros respetivamente, e que ambos encontraram C erros em comum. Agora suponhamos que o primeiro examinador tem a probabilidades de detetar um erro ao passo que o outro tem b probabilidades de encontrar um erro. Se o número total de erros tipográficos na tese for T , então $A = aT$ e $B = bT$.

Mas se os dois examinadores estiverem a rever a tese *independentemente*, também conhecemos outro facto chave, o de que $C = abT$. Assim, $AB = abT^2 = CT$ e, conseqüentemente, o número total de erros é $T = AB/C$, independentemente dos valores de a e b . Uma vez que o número total de erros encontrados pelos examinadores (notando que não devemos contar duas vezes os C erros que ambos encontraram) é $A + B - C$, o que significa que o número total de erros que eles não detetaram é apenas $T - (A + B - C)$ que por sua vez é $(A - C)(B - C)/C$. Por outras palavras, é o produto do número que cada um deles encontrou e que o outro não encontrou, dividido pelo número de erros que ambos encontraram. Faz sentido. Se ambos tivessem encontrado muitos erros, mas não tivessem encontrado erros em comum, não seriam muito bons revisores e era provável que houvesse muitos mais erros que não tivessem sido encontrados por nenhum dos dois. Na minha tese, tínhamos $A = 32$, $B = 23$ e $C = 16$, por isso, espera-se que o número de erros não detetados seja $(16 \times 7)/16 = 7$.

Este tipo de argumentos pode ser usado em muitas situações. Suponhamos que diferentes prospectores procuram independentemente lençóis de petróleo: quantos poderão ficar por encontrar? Ou se um grupo de ecologistas quiser saber quantas espécies de mamíferos ou de aves podem existir numa região florestal com vários observadores a fazerem um censo com uma duração de 24 horas.

Um tipo semelhante de problema surgiu na análise literária. Em 1976, dois estatísticos usaram a mesma abordagem para calcular a extensão do léxico de William Shakespeare, investigando o número de palavras diferentes que usava nas suas obras e tendo em conta diversos usos. Shakespeare escreveu aproximadamente um total de 900 000 palavras. Destas, 31 534 são diferentes entre si e 14 376 surgem apenas uma vez. Há 4343 palavras que surgem apenas duas vezes e 2292 surgem apenas três vezes. Prevvia-se que Shakespeare conhecia pelo menos 35 000 vocábulos que não foram usados nas suas obras. Isto significa que ele teria provavelmente um léxico de cerca de 66 500 vocábulos. Surpreendentemente, é mais ou menos o mesmo número de vocábulos que o leitor domina.

5 | RÂGUEBI E RELATIVIDADE

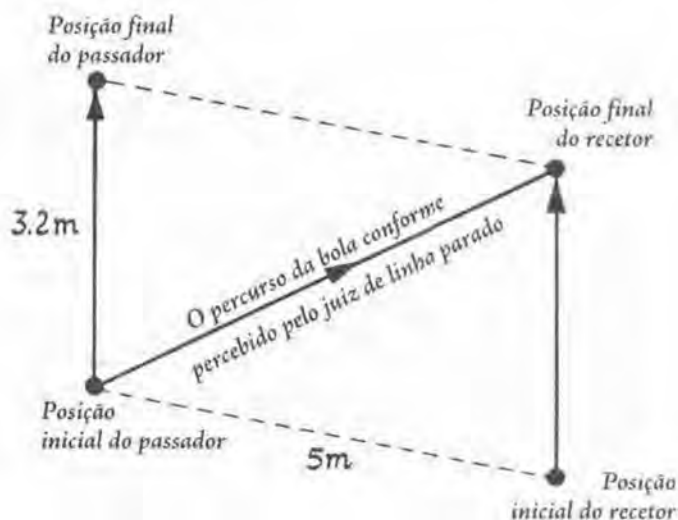
O râguebi é um jogo que não me posso orgulhar de conhecer ao pormenor, se é que me entendem. Consigo compreender os princípios gerais, claro. Quero dizer, sei que o objectivo geral é levar a bola para o outro lado do campo e depositá-la para lá da linha na outra ponta do campo e que, para cumprir este objetivo, ambas as equipas têm permissão para imprimir uma certa violência aos seus atos e fazer coisas aos seus semelhantes que, noutro desporto, resultaria em catorze dias de suspensão e uma forte reprimenda do árbitro.

P. G. Wodehouse, *Very Good, Jeeves*

A relatividade do movimento não tem de ser um problema apenas para Einstein. Quem é que nunca passou pela experiência de estar sentado numa carruagem de comboio parada na estação e ter subitamente a sensação de estar em movimento, apenas para reconhecer que foi o comboio na linha paralela que avançou, sem que o comboio onde se encontrava se tivesse movido?

Eis outro exemplo: há cinco anos, passei duas semanas a visitar a universidade de New Wales em Sidney, numa altura em que a Taça do Mundo de Râguebi estava a dominar as notícias e o interesse do público. Depois de ver vários jogos na televisão, reparei num interessante problema de relatividade que passou despercebido às celebridades no estúdio. Um passe para diante é relativo a quê? As regras são claras: um passe para diante (também referido como *en avant*) ocorre quando a bola é atirada na direção da linha de gol oposta. Mas quando os jogadores estão em movimento, a situação torna-se mais difícil de avaliar para um observador, devido à relatividade do movimento.

Imagine que dois jogadores atacantes estão a correr (pela página acima) em linhas retas paralelas separadas por uma distância de 5 metros a uma velocidade de 8 metros por segundo em direção à linha do adversário. Um dos jogadores, o "recetor", está 1 metro atrás do outro, o "passador",



que tem a bola. O passador atira a bola a uma velocidade de 10 metros por segundo na direção do recetor. A velocidade da bola em relação ao chão é, na realidade, de $\sqrt{(10^2 + 8^2)} = 12,8$ metros por segundo e demora 0,4 segundos a percorrer os 5 metros que separam os jogadores. Durante este intervalo, o recetor correu uma distância superior a $8 \times 0,4 = 3,2$ metros. Quando o passe foi lançado, ele estava 1 metro atrás do passador, mas quando apanha a bola está 2,2 metros à frente dele do ponto de vista de um juiz de linha que esteja ao nível do passe original. O juiz acredita que houve um passe para diante e faz sinal com a bandeira. No entanto, o árbitro está a correr ao longo da jogada e não vê a bola a avançar, por isso deixa o jogo continuar.

6 | NUMA RODA-VIVA

O meu coração é como uma roda.

Paul McCartney, "Let Me Roll It"

Certo fim de semana, notei que os jornais estavam a falar de propostas para introduzir um limite de velocidade mais restritivo de vinte milhas por hora nas áreas residenciais no Reino Unido e para usar câmaras com radar para auxiliar na aplicação desse limite, quando fosse possível. Esquecendo as questões relativas à segurança no trânsito, há alguns aspetos do movimento rotacional que sugerem que as câmaras com radar podem acabar por apanhar uma grande quantidade de ciclistas perplexos, aparentemente a exceder o limite de velocidade por fatores bastante significativos. Como é que isto é possível?

Suponhamos que uma bicicleta está a deslocar-se a uma velocidade V em direção a um detetor de velocidade. Isto significa que tanto o eixo da roda como o corpo do ciclista estão a deslocar-se a uma velocidade V em relação ao chão. Mas observemos com mais atenção o que está a acontecer em pontos diferentes da roda em movimento rotativo. Se a roda não derapar, a velocidade da extremidade da roda que está em contacto com o chão deverá ser zero. Se a roda tiver um raio R e estiver a girar com uma velocidade angular constante de Ω voltas por segundo, a velocidade do ponto de contacto também pode ser descrita como $V - R\Omega$. Este valor tem de ser igual a zero, pelo que $V = R\Omega$. A velocidade de deslocação do centro da roda é V , mas a velocidade de deslocação da parte superior da roda é a soma de V e da velocidade rotacional, que é igual a $V + R\Omega$, sendo, portanto, igual a $2V$. Se uma câmara determinar a velocidade de uma bicicleta a aproximar-se ou a afastar-se medindo a velocidade da parte superior da roda, registará uma velocidade correspondente ao dobro da

velocidade a que o ciclista está a deslocar-se. Este problema pode ser interessante para os meus sábios colegas, mas recomendo que usem um bom par de guarda-lamas.

7 | UMA NOÇÃO DE PROPORÇÃO

Só pode encontrar a verdade através da lógica quem já tiver encontrado a verdade sem ela.

G. K. Chesterton

A medida que uma pessoa cresce, fica mais forte. Vemos todo o tipo de exemplos de aumento da força correspondente ao aumento de tamanho no mundo que nos rodeia. A força superior dos lutadores de boxe, praticantes de luta livre e halterofilistas mais pesados é verificável pela necessidade de escalonar as competições de acordo com o peso dos participantes. Mas a que velocidade aumenta a força em relação ao aumento de peso ou de tamanho? Consegue acompanhar o ritmo do crescimento? Afinal de contas, um gatinho bebé consegue manter a cauda erguida na vertical, ao passo que a sua mãe, muito maior, não consegue – a cauda curva-se sob o seu próprio peso.

Exemplos simples podem ser muito esclarecedores. Pegue num *gressino* comprido e parta-o ao meio. Agora faça o mesmo com outro muito mais comprido. Se o agarrar sempre à mesma distância do ponto de quebra, descobrirá que não é mais difícil partir o *gressino* comprido do que o curto. Uma breve reflexão ajuda-nos a perceber o motivo. O *gressino* parte-se numa secção transversal ao longo do seu comprimento. É ali que se dá toda a ação: uma pequena linha de ligações moleculares do *gressino* quebra-se e ele parte-se. O resto do *gressino* é irrelevante. Mesmo que tivesse 100 metros de comprimento, não teria mais dificuldade em partir aquela fina secção de ligações num ponto ao longo do seu comprimento. A força do *gressino* é determinada pelo número de ligações moleculares que têm de ser quebradas na secção transversal. Quanto maior for essa área, mais ligações terão de ser que-

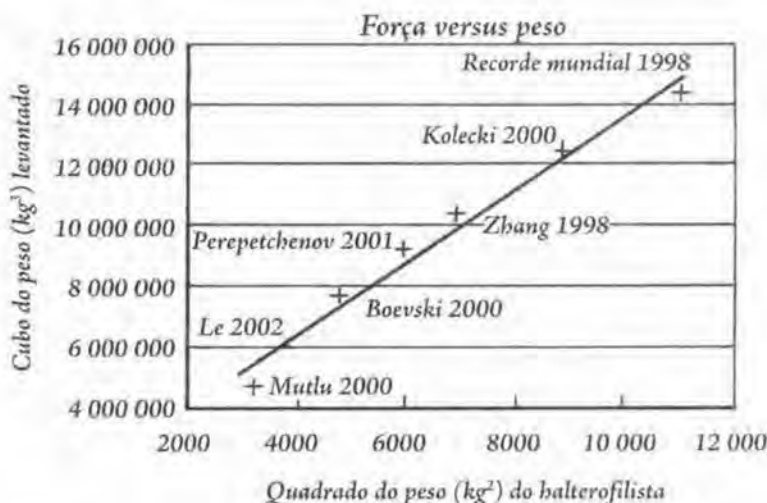
bradas e mais forte será o *gressino*. Assim, a força é proporcional à secção transversal que, por sua vez, é proporcional a uma determinada medida do seu diâmetro ao quadrado.

Coisas quotidianas como *gressinos* e halterofilistas têm uma densidade constante que é simplesmente determinada pela densidade média dos átomos que os compõem. Mas a densidade é proporcional à massa dividida pelo volume, que é a massa dividida pelo tamanho do cubo. Aqui na superfície da Terra, a massa é proporcional ao peso e, por isso, observamos a mesma "lei" da proporcionalidade que se aplica aos objetos relativamente esféricos:

$$(\text{Força})^3 \propto (\text{peso})^2$$

Esta simples regra permite-nos compreender todo o tipo de coisas. A relação entre força e peso diminui à medida que $\text{Força/Peso} \propto (\text{peso})^{-1/3} \propto 1/(\text{tamanho})$. Portanto, à medida que uma pessoa cresce, a sua força não acompanha o aumento de peso. Mesmo que todas as dimensões se expandissem uniformemente em tamanho, acabaria por ficar com um peso demasiado grande para ser suportada pelos seus ossos e partir-se-ia. É por isso que existe um tamanho máximo para estruturas terrestres compostas por átomos e moléculas, sejam elas dinossauros, árvores ou edifícios. Se aumentarmos as dimensões destas em forma e tamanho, acabarão por ficar tão grandes que o peso quebrará as ligações moleculares na base e abaterão sob o seu próprio peso.

Começámos por mencionar alguns desportos em que a vantagem em termos de tamanho e peso é tão significativa que os atletas são divididos em categorias de acordo com o seu peso. A nossa "lei" diz que devemos esperar uma correlação em linha reta quando calculamos o cubo do peso levantado em relação ao quadrado do peso dos halterofilistas. É isto que acontece quando se traça esse gráfico para os recordes mundiais masculinos em função das categorias de peso:



A correspondência é quase exata! Às vezes a matemática pode simplificar a nossa vida. O halterofilista no ponto mais elevado da linha traçada segundo esta “lei” é o halterofilista que consegue levantar o maior peso por quilo, ao passo que o halterofilista mais pesado, o que levanta o maior peso global, é na verdade relativamente o mais fraco se tivermos em conta quanto pesa.

8 | POR QUE É QUE A OUTRA FILA AVANÇA SEMPRE MAIS DEPRESSA?

A relva do jardim do vizinho é sempre mais verde.
O sol brilha sempre mais do outro lado.

Cantado por Petula Clark

Já deve ter reparado que sempre que entra numa fila no aeroporto ou nos correios, as outras filas parecem estar a avançar mais depressa. Quando o trânsito está parado na autoestrada, as outras faixas avançam sempre mais depressa do que a sua. Mesmo que troque de fila, a sua continua a ser a mais lenta. Esta situação é frequentemente referida como “Lei de Sod” e parece ser uma manifestação de um princípio profundamente antagónico no centro da realidade. Ou talvez seja apenas uma manifestação da paranoia humana ou um registo seletivo de evidências. Ficamos impressionados com as coincidências sem pararmos para refletir sobre situações muito mais numerosas em que a coincidência não se deu e que nós simplesmente não registámos essa informação. De facto, a razão pela qual parece estar tantas vezes na fila mais lenta pode não ser uma ilusão. É uma consequência do facto de *estar* habitualmente na fila mais lenta.

A razão é simples. Em média, as filas e faixas mais lentas são as que têm mais pessoas e veículos. Consequentemente, é mais provável que esteja numa destas e não numa das filas ou faixas mais rápidas que têm menos pessoas.

A condição “em média” é muito importante neste contexto. Cada fila em particular terá características distintas – pessoas que se esquecem da carteira, que têm um carro que não anda a mais de 50 quilómetros por hora, etc. O leitor não estará invariavelmente na fila mais lenta, mas

em média, tendo em conta todas as filas em que se coloca, é mais provável que tenha estado nas filas com mais pessoas.

Este tipo de autosseleção é uma inclinação que pode ter consequências significativas para a ciência e para a análise de dados, especialmente se não for tida em conta. Suponhamos que quer determinar se as pessoas que vão regularmente à igreja são mais saudáveis do que as que não o fazem. É um erro que deve evitar. As pessoas menos saudáveis estão impossibilitadas de ir à igreja, portanto, se se limitar a contar as pessoas na igreja e a observar o seu estado de saúde obterá um resultado falso. Da mesma forma, se considerarmos o Universo, podemos ter em mente um "princípio", inspirado por Copérnico, que nos diz que não devemos pensar que a nossa localização no Universo é especial. Contudo, embora não devamos esperar que a nossa localização seja especial *em todos os sentidos*, seria um erro muito grave considerar que ela não é especial *em nenhum* sentido. Só pode surgir vida em locais onde existam circunstâncias especiais: é mais provável que possa ser encontrada vida onde existam estrelas e planetas. Estas estruturas formam-se em locais especiais, onde a abundância de poeiras interestelares é superior à média. Desta forma, quando estudamos ciências ou quando nos confrontamos com dados, o que é mais importante averiguar relativamente aos dados é se está presente alguma inclinação que nos leve a tirar preferencialmente uma conclusão e não outra do conjunto de evidências disponível.

9 | DOIS É BOM, TRÊS É DE MAIS

Tudo o que sobe, desce.

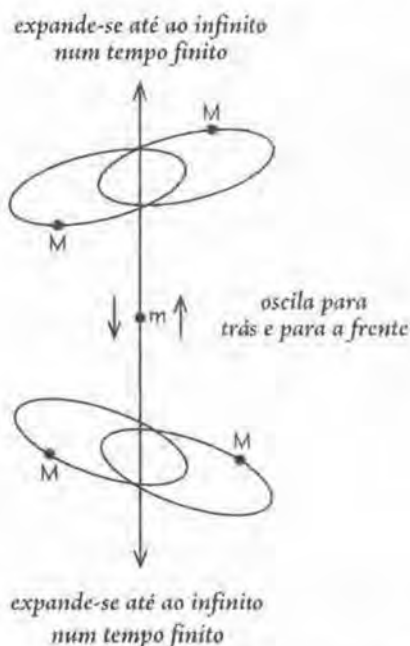
Anónimo

Duas pessoas que se dão muito bem podem ver a sua relação a ser destabilizada pela chegada de uma terceira pessoa ao grupo. Este facto é ainda mais aparente quando a gravidade é a força de atração envolvida. Newton ensinou-nos que duas massas podem permanecer numa órbita estável em torno do centro da sua massa sob as suas forças gravitacionais mútuas – como acontece com a Terra e a Lua. Mas se um terceiro corpo com uma massa semelhante for introduzido no sistema, de um modo geral acontece uma mudança muito dramática. Um dos corpos é expulso do sistema pelas forças gravitacionais, ao passo que os dois corpos restantes são atraídos para uma órbita mais estável e ficam mais fortemente ligados.

Este processo “fisga” simples é a fonte de uma propriedade contraintuitiva da teoria da gravidade de Newton, descoberta por Jeff Xia em 1992. Em primeiro lugar, consideremos quatro partículas de massa igual M e disponhamo-las em dois pares a orbitar em dois planos paralelos e com direções rotacionais opostas, para não haver rotação global do sistema. Agora, introduzamos uma quinta partícula muito mais leve, m , que oscila para trás e para a frente na perpendicular em relação aos centros da massa dos dois pares. O grupo de cinco partículas irá expandir-se para um tamanho infinito num período de tempo finito!

Como é que isto acontece? A pequena partícula oscilante corre de um par para o outro e ao alcançar o outro par, cria um pequeno problema com três corpos e é repelida. O par recua para conservar o impulso. Em seguida, a partícula mais leve desloca-se para junto do outro par e repete-se o mesmo cenário. Isto acontece repetidamente, infinitamente, e acelera tanto os dois

pares que alcançam uma separação infinita num período de tempo finito, sofrendo um número infinito de oscilações no processo.



Este exemplo resolve um velho problema colocado pelos filósofos acerca da possibilidade de desempenhar um número infinito de ações num período de tempo finito. Claramente, num mundo newtoniano em que não existe um limite de velocidade, isso é possível. Infelizmente (ou talvez felizmente), este comportamento não é possível se tivermos em conta a relatividade de Einstein. Não é possível transmitir informações mais depressa do que a velocidade da luz e as forças gravitacionais não podem tornar-se arbitrariamente fortes na teoria do movimento e gravitação de Einstein. Da mesma forma, as massas não podem aproximar-se e afastar-se arbitrariamente. Quando duas massas M se aproximam a uma distância inferior a $4GM/c^2$, em que G é a constante gravitacional de Newton e c é a velocidade da luz, forma-se em volta delas um "horizonte" de não-retorno e elas formam um buraco negro do qual não podem sair.

O efeito "fisga" da gravidade pode ser demonstrado no seu jardim através de uma experiência simples. Esta mostrará como três corpos podem combinar-se para criar uma grande retração enquanto tentam conservar

o impulso ao passarem perto uns dos outros (no caso dos corpos astronômicos) ou ao colidirem (como acontecerá na nossa experiência).

Estes três corpos serão a Terra, representada por uma bola grande (como uma bola de basquetebol ou uma bola de futebol de superfície lisa) e uma bola pequena (como uma bola de pingue-pongue ou de tênis). Segure nas duas bolas, aproximadamente à altura do seu peito, mantendo a bola pequena ligeiramente acima da bola grande, e depois deixe-as cair ao chão ao mesmo tempo. A bola grande tocará primeiro no chão e ressaltará, batendo na bola pequena enquanto esta ainda está a descer. O resultado é bastante dramático. A bola pequena sobe a uma altura *nove* vezes superior à que teria subido se tivesse sido apenas deixada cair da mesma altura. Talvez seja melhor não fazer esta experiência dentro de casa!

* A bola de basquetebol ressaltará do chão a uma velocidade V e atinge a bola de pingue-pongue quando esta ainda está a cair à velocidade V . Por tanto, em relação à bola de basquetebol, a bola de pingue-pongue eleva-se à velocidade $2V$ depois de a sua velocidade ter sido invertida pela colisão. Tendo em conta que a bola de basquetebol está a deslocar-se à velocidade V em relação ao chão, isto significa que a bola de pingue-pongue está a elevar-se à velocidade $2V + V = 3V$ em relação ao chão após a colisão. Considerando que a altura alcançada é proporcional a V^2 , isto significa que a bola irá elevar-se a uma altura $3^2 = 9$ vezes superior àquela a que se elevaria sem a colisão com a bola de basquetebol. Na prática, a perda de energia sofrida nos ressaltos fará a bola subir a uma altura ligeiramente inferior.

10 | O MUNDO É MESMO PEQUENO

O mundo é pequeno, mas corremos em grandes círculos.

Sashia Azevedo

Quantas pessoas conhece o leitor? Tomemos como exemplo médio um número redondo, como 100. Se cada um dos seus conhecidos conhecer outras 100 pessoas diferentes, um único passo liga o leitor a 10 000 pessoas: o que mostra que é uma pessoa muito mais bem relacionada do que julga. Ao fim de n passos como este, o leitor estará ligado a $10^{2(n+1)}$ pessoas. No mês em que este texto foi escrito estimava-se que a população mundial fosse de 6,65 mil milhões, ou seja, $10^{9,8}$ o que significa que se $2(n+1)$ for maior que 9,8, o leitor tem um número de ligações superior à população mundial. Isto acontece quando n é superior a 3,9 pelo que bastarão 4 passos para a situação se verificar.

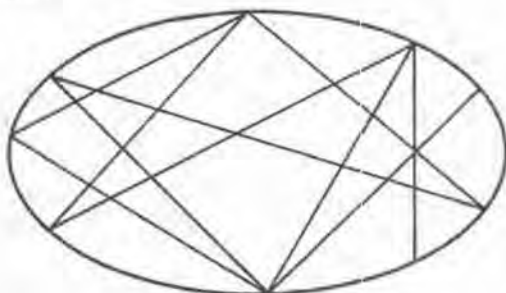
Esta conclusão é bastante surpreendente. Parte de muitas suposições simples que não correspondem à realidade, começando pela suposição de que os amigos dos seus amigos são todos diferentes entre si. Mas se repetirmos mais cuidadosamente a contagem de modo a ter em conta este facto, a diferença não é grande. Bastam 6 passos para o ligar a qualquer habitante deste planeta. Experimente e ficará surpreendido com a frequência com que menos de 6 passos o ligam a pessoas famosas.

Há aqui uma suposição oculta que não está bem testada se considerarmos a sua ligação ao Primeiro-Ministro, ao David Beckham ou ao Papa. Estará surpreendentemente próximo destas pessoas famosas porque elas estão ligadas a muitas outras pessoas. Mas experimente estabelecer uma ligação com um membro de uma tribo indígena da Amazónia ou com um pastor da Mongólia e descobrirá que a cadeia que o liga a estas pessoas tem muitos mais elos. Poderá até não conseguir estabelecer a ligação. Estes

indivíduos vivem em grupos fechados que têm poucas ligações com o mundo exterior.

X-----X-----X-----X-----X

Se a sua rede de relações for uma cadeia ou um círculo, só estará ligado por uma pessoa em cada um dos extremos e o nível geral de ligações é baixo. No entanto, se estiver num círculo de relações à qual são adicionadas aleatoriamente outras ligações conseguirá ir de uma ponta à outra muito rapidamente.



Em anos recentes, temos vindo a perceber os efeitos dramáticos de poucas ligações de longo alcance na conectividade geral. Um conjunto de núcleos que produzem muitas ligações a locais próximos pode ser eficazmente ligado entre si se adicionarmos entre eles algumas ligações de longo alcance.

Estas noções são importantes quando se trata de determinar a cobertura necessária para ligar todos os utilizadores através de uma rede de comunicações móveis, ou de que forma alguns indivíduos infectados por um vírus podem espalhar uma doença através das suas interações com os membros de uma população. Quando as companhias aéreas tentam planear os nós de ligação e os percursos de forma a minimizar a duração das viagens ou a maximizar o número de cidades que conseguem ligar com uma única escala ou a minimizar os custos, precisam de compreender as propriedades inesperadas destas redes de "pequenos mundos".

O estudo das nossas ligações mostra-nos que o "mundo" existe em muitos níveis: as ligações de transportes, as linhas telefónicas e os caminhos de e-mail criam redes de interligações que nos unem de formas inesperadas. Tudo está mais perto do que pensávamos.

11 | ESTABELECECER PONTES

Como uma ponte sobre águas turbulentas.

Paul Simon e Art Garfunkel,

"Bridge Over Troubled Water"

Um dos grandes feitos da engenharia foi a construção de pontes que nos permitem atravessar rios e ravinas que seriam, de outra forma, intransponíveis. Estes vastos projetos de construção têm frequentemente uma qualidade estética que os coloca entre as maravilhas do mundo moderno. A elegante ponte Golden Gate em São Francisco, a notável Clifton Suspension Bridge de Brunel e a Ponte Hercílio Luz no Brasil têm formas espetaculares que partilham um aspeto regular e semelhante. Que formas são essas?

Existem duas formas interessantes que se manifestam quando pesos e correntes são suspensos, e que são vulgarmente confundidas ou erradamente consideradas iguais. O problema mais antigo era o de descrever a forma assumida por uma corrente ou corda suspensa cujas extremidades estão fixas em dois pontos no mesmo nível horizontal. O leitor pode verificar facilmente este facto. A primeira pessoa a afirmar que sabia qual era essa forma foi Galileu que, em 1638, declarou que uma corrente suspensa nestas circunstâncias e sujeita à gravidade assumiria a forma de uma parábola (expressada como $y^2 = Ax$ em que A representa qualquer número positivo). Mas em 1669, Joachim Jungius, um matemático alemão que tinha um interesse especial nas aplicações da matemática a problemas físicos, provou que Galileu estava errado. A determinação da equação para a corrente suspensa acabou por ser calculada por Gottfried Leibniz, Christiaan Huygens, David Gregory e Johann Bernoulli em 1691, um ano depois de Johann Bernoulli ter anunciado publicamente este problema como um desafio. Esta curva foi inicialmente designada por *catenária* por Huygens

numa carta a Leibniz, uma palavra que deriva da palavra latina *catena*, que significa *cadeia*. Na língua inglesa, o equivalente anglicizado “catenary” parece ter sido introduzido pelo Presidente Thomas Jefferson numa carta a Thomas Paine, datada de 15 de setembro de 1788, relativamente ao desenho de uma ponte. A forma também tem sido referida como *chainette* ou curva *funicular*.

A forma de uma catenária reflete o facto de a sua tensão suportar o peso da própria corrente e o peso total suportado por qualquer ponto desta é, assim, proporcional ao comprimento total da corrente entre esse ponto e o ponto mais baixo da corrente. A equação relativa à corrente suspensa tem a forma $y = B \cosh(x/B)$ em que B é a tensão constante da corrente dividida pelo seu peso por unidade de comprimento¹. Se segurar dois extremos de uma corrente suspensa e os aproximar ou afastar, a forma da corrente continua a ser descrita por esta fórmula, mas com um valor diferente de B para cada posição. Esta curva também pode ser encontrada quando se tenta descobrir a forma que coloca o centro de gravidade da corrente no ponto mais baixo possível.



Gateway Arch em St. Louis

Outro exemplo espetacular de uma catenária construída pelo Homem pode ser visto em St. Louis, Missouri, cujo Gateway Arch é uma catenária invertida. Esta é a forma ideal para um arco autossustentável, que minimiza o cisalhamento porque a tensão é sempre dirigida ao longo da linha do arco em direção ao chão. A sua fórmula matemática exata está inscrita no arco. Por estes motivos os arcos catenários são frequentemente usados

pelos arquitetos para otimizar a força e estabilidade das estruturas; um exemplo notável pode ser encontrado nos arcos incrivelmente altos da Igreja da Sagrada Família de Antoni Gaudí em Barcelona.

Outro belo exemplo é o Museu da Artilharia desenhado por John Nash em 1819, que fica situado nos limites do parque Woolwich Common, em Londres. O seu telhado peculiar, semelhante a uma tenda, inspirado pela forma das tendas dos soldados, tem a forma de metade de uma curva catenária.



Edifício do Museu da Artilharia desenhado por John Nash

No entanto, existe uma grande diferença entre uma corrente suspensa e uma ponte suspensa como a ponte Clifton ou a Golden Gate. As pontes suspensas não têm apenas de suportar o peso dos seus próprios cabos ou correntes. A maior parte do peso a suportar é o tabuleiro da ponte. Se o tabuleiro for horizontal e tiver uma área transversal de densidade constante a todo o comprimento, a equação para determinar a forma do cabo de suporte passa a ser uma parábola $y = x^2/2B$, em que B é (tal como acontecia no exemplo da corrente suspensa) uma constante igual à tensão dividida pelo peso por unidade de comprimento do tabuleiro da ponte.

Um dos exemplos mais notáveis é a Clifton Suspension Bridge em Bristol, desenhada por Isambard Kingdom Brunel em 1829, mas cuja construção só foi concluída em 1865, três anos após a sua morte. A sua bela forma parabólica continua a ser um monumento às grandes obras de engenharia desde Arquimedes.

Por que é que as crianças já não colecionam coisas? O que é que aconteceu àquelas coleções de selos meticulosamente mantidas...?

Woman's Hour, BBC Radio 4

No fim de semana passado, escondidos entre alguns livros na parte de trás da minha estante, encontrei dois conjuntos de cartas que eu tinha colecionado em criança. Cada um deles continha 50 fotografias a cores de carros clássicos, impressas em alta qualidade, e tinham no verso uma descrição bastante detalhada do *design* e das especificações mecânicas. As coleções de cartas deste tipo foram em tempos uma grande moda. Havia coleções sobre aviões de guerra, animais, flores, navios e desportistas – uma vez que todas elas pareciam ser dirigidas a um público do sexo masculino – que podiam ser adquiridas comprando muitas embalagens de pastilhas elásticas, cereais de pequeno-almoço e embalagens de chá em saquetas. Das que tinham uma temática desportiva, à semelhança dos autocolantes que saem agora nos bolos pré-embalados e nos pacotes de batatas fritas, o jogo preferido de todos era o futebol (nos EUA era o beisebol) e eu sempre duvidei que todas as cartas fossem produzidas em quantidades iguais. Por algum motivo, todas as pessoas pareciam acabar sempre por ter em falta a carta do “Bobby Charlton” que era necessária para completar a coleção. Todas as outras cartas podiam ser facilmente adquiridas trocando as repetidas com os nossos amigos, mas esta carta essencial parecia faltar a todos.

Foi com alívio que descobri que os meus filhos se dedicam a um colecionismo semelhante. Os objetos colecionados podem ter mudado, mas a ideia básica é a mesma. Então onde é que a matemática entra nesta história? A pergunta interessante neste cenário é quantas cartas podemos esperar comprar para completar a coleção, se partirmos do pressuposto que todas são produzidas em números iguais, tendo assim igual probabi-

lidade de ser encontradas na próxima embalagem que abrirmos. Os conjuntos de cartas sobre carros que encontrei tinham 50 cartas cada um. A primeira carta que vou encontrar vai sempre ser nova, mas então e a segunda? Há uma probabilidade de 49/50 de não a ter. A seguinte terá uma probabilidade de 48/50 de ser nova, e por aí em diante.

Depois de ter adquirido 40 cartas, haverá uma probabilidade de 10/50 de que a próxima carta adquirida seja nova. Assim, uma pessoa terá de comprar em média outras 50/10 ou mais 5 cartas para ter uma maior probabilidade de conseguir mais uma carta que ainda não tenha para completar a coleção. Desta forma, o número total de cartas que, em média, precisará de comprar para completar a coleção de 50 será a soma de 50 termos:

$$50/50 + 50/49 + 50/48 + \dots + 50/3 + 50/2 + 50/1$$

em que o primeiro termo é o caso certo da primeira carta que obtém e cada termo sucessivo diz-lhe quantas cartas adicionais precisa de comprar para conseguir cada uma das restantes cartas em falta no conjunto de 50.

Como podem existir coleções com diferentes números de cartas, consideremos a aquisição de um conjunto com um número indeterminado de cartas, número esse que designaremos por N . A mesma lógica diz-nos que, em média, teremos de comprar um total de:

$$(N/N) + (N/(N-1)) + (N/(N-2)) + \dots + N/2 + N/1 \text{ cartas}$$

Se excluirmos o fator comum N dos numeradores de cada termo, será apenas:

$$N(1 + 1/2 + 1/3 + \dots + 1/N)$$

A soma dos termos entre parênteses é a famosa série "harmônica". Quando N se torna grande, adquire um valor aproximado de $0,58 + \ln(N)$ em que $\ln(N)$ é o logaritmo natural de N . Assim, quando N se torna realisticamente grande, observamos que o número médio de cartas que precisamos de comprar para completar o nosso conjunto é aproximadamente:

$$\text{Cartas necessárias} \approx N \times [0,58 + \ln(N)]$$

Para os meus conjuntos de 50 cartas, a resposta é 224,5, e deveria ter esperado comprar em média 225 cartas para completar o conjunto de 50. Curiosamente, os nossos cálculos mostram que se torna muito mais difícil completar a segunda metade da coleção do que a primeira. O número de cartas que é preciso comprar para colecionar $N/2$ cartas para meia coleção é:

$$(N/N) + (N/(N-1)) + (N/(N-2)) + \dots + N/((\frac{1}{2}N+1))$$

que é a diferença entre N vezes a série harmónica somada a N e somada a $N/2$ termos, pelo que:

Cartas necessárias para meia coleção $\approx N \times [\ln(N) + 0,58 - \ln(N/2) - 0,58] = N \ln(2) = 0,7N$

ou apenas 35 para completar a primeira metade da minha coleção de 50.

Pergunto-me se os fabricantes fizeram estes cálculos. Devem tê-los feito porque eles são o que permite calcular o lucro máximo que se pode esperar obter a longo prazo com a divulgação de uma coleção de cartas específica. É provável que seja apenas o lucro máximo possível, porque os colecionadores fazem trocas de cartas, o que lhes permite adquirir novas cartas sem terem de as comprar.

Que impacto podem ter as trocas de cartas com os seus amigos?

Suponhamos que tem A amigos e que todos contribuem com cartas para fazer $A + 1$ coleções, de modo a cada um de vocês ter uma. Quantas cartas serão necessárias para o fazer? Em média, quando o número de cartas N é grande e é feita troca de cartas, a resposta aproxima-se de:

$$N \times [\ln(N) + A \ln(\ln N) + 0,58]$$

Por outro lado, se cada um de vocês tiver completado uma coleção sem fazer trocas, terão precisado de aproximadamente $(A+1)N[\ln(N) + 0,58]$ cartas para completar $A + 1$ coleções separadas. Para $N = 50$, o número de compras de cartas poupadas seria $156A$. Mesmo que $A = 1$, é uma poupança considerável.

Se tiver alguns conhecimentos de estatística, poderá querer demonstrar que o desvio que pode ser esperado no resultado $N \times [0,58 + \ln(N)]$

é aproximadamente de $1,3N$. Na prática, é um desvio muito significativo, porque significa que tem 66% de probabilidades de precisar de colecionar mais ou menos $1,3N$ acima da média. Para a nossa coleção de 50 cartas, esta incerteza no número de compras esperadas é de 65. Há alguns anos contava-se a história de um consórcio que estava a analisar as lotarias para calcular o número médio de bilhetes que era necessário comprar para ter uma melhor probabilidade de acertar em todos os números possíveis – incluindo o número vencedor. Os membros deste consórcio esqueceram-se de incluir a variação provável do resultado médio, mas tiveram a sorte de descobrir que entre os milhões de bilhetes que tinham comprado se contava o bilhete premiado.

Se a probabilidade de cada carta aparecer não for a mesma, o problema torna-se mais complexo, mas ainda assim solúvel. Nesse caso é mais semelhante ao problema das coleções de moedas em que um colecionador tenta obter moedas com todas as datas. Não sabe se foram cunhadas em número igual em todos os anos (embora muito provavelmente não tenham sido), nem quantas foram posteriormente retiradas de circulação, pelo que não pode contar com uma probabilidade igual de encontrar um *penny* de 1840 e um de 1890. Mas se por acaso o leitor encontrar um *penny* inglês de 1933 (ano em que foram cunhados apenas 7 e só é conhecido o paradeiro de 6), avise-me.

13 | DA NUMERAÇÃO ROMANA AO CRÍQUETE

É uma ideia profundamente errada... a de que devemos cultivar o hábito de pensar no que estamos a fazer. A verdade é exatamente o contrário. A civilização avança com o aumento da quantidade de operações importantes que conseguimos efetuar sem pensar.

Alfred North Whitehead

O prisioneiro desesperado é encarcerado numa cela escura, húmida e esquecida. Os dias, meses e anos passam lentamente. E seguir-se-lhes-ão muitos mais. Conhecemos esta cena dos filmes. No entanto, habitualmente há um interessante subtexto matemático. O prisioneiro tem vindo a contar os dias por meio de sistemas de marcas na parede. Os artefactos humanos mais antigos conhecidos com registos de contagens na Europa e África remontam a mais de 30 000 anos e mostram elaborados grupos semelhantes de marcas a seguir os dias do mês e a acompanhar as fases da Lua.

O padrão europeu de marcação por barras recupera o método de contagem pelos dedos e resulta num conjunto de 4 linhas (mais ou menos) verticais |||| que é finalizado por uma linha diagonal para fechar um conjunto de cinco. Passa para o seguinte conjunto de cinco ||||| e por aí fora. O conjunto de quatro barras verticais fechado com uma barra diagonal denotando cada item contado mostra um método de contagem anterior aos nossos sistemas formais e conduziu à adoção dos numerais romanos I, II e III ou às varetas do sistema numérico chinês básico. Estão estreitamente ligados a métodos simples de contagem pelos dedos e usam grupos de 5 e de 10 como base para os conjuntos de marcas acumuladas. Os sistemas antigos registavam as marcas com entalhes feitos em osso ou numa vara de madeira. Os primeiros números eram marcados por entalhes

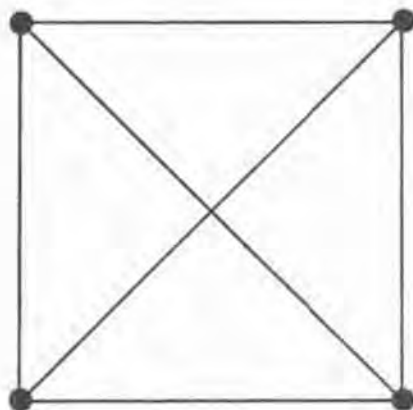
individuais, mas a meia cruz, com a forma de um V, começou a ser usada para marcar o número 5 e uma cruz completa, X, começou a designar o número 10, dando assim origem aos numerais romanos V e X; o quatro era formado pela soma, como IIII, ou pela subtração, como IV. O sistema de registo de contagens continuou a ser uma importante atividade oficial em Inglaterra até 1826, com o Tesouro a usar grandes varas marcadas (à semelhança dos metros de alfaiate) para registar as grandes somas que entravam e saíam dos seus cofres. Este uso também foi a fonte dos vários significados da palavra inglesa *score*, que tanto pode significar fazer uma marca, ou registar, como pode designar uma quantidade de 20. Quando havia um valor em dívida ao Tesouro, a vara com as marcas era dividida ao meio e era dada metade ao devedor, ficando a outra metade para o Tesouro. Quando a dívida era paga, as duas metades eram unidas para confirmar que correspondiam ao total.

Contar estas marcas é um trabalho laborioso, especialmente se forem obtidos grandes totais. É necessário contar mentalmente as marcas individuais e depois também é necessário somar o número de conjuntos. Na América Latina, encontra-se ocasionalmente o uso de um sistema que utiliza uma acumulação gradual dos números traçando quatro linhas que formam um quadrado, completado por duas linhas diagonais cruzadas no seu interior .

Estamos familiarizados com uma variante desta forma de registo quando contabilizamos as pontuações num jogo de críquete, fazendo seis marcas em duas filas e dois traços horizontais, um "ponto" se não tiver sido marcado nenhum *run*, o número marcado, ou um "w" se um *wicket* falhar. Outros símbolos denotam *wides*, *byes*, *no-balls* e *leg byes*. Se não forem marcados *runs* em nenhuma das seis bolas, os seis pontos são unidos de modo formar um W, para indicar um *wicket maiden*. Desta forma, basta olhar para uma tabela de pontuações para conhecer de imediato o estado do jogo e o número de *runless overs*.

É tentador combinar os sistemas dos jogos Sul-americanos com os da pontuação de um jogo de críquete para criar um auxiliar de memória ideal para qualquer pessoa que queira utilizar o sistema decimal e não esteja disposta a contar individualmente cada uma das dez marcas verticais para calcular o total. Primeiro conta-se de um a quatro, colocando quatro pontos nos cantos de um quadrado, depois conta-se de cinco a oito,

traçando os lados do quadrado, e depois conta-se o nove e o dez, acrescentando as duas linhas diagonais no interior do quadrado. Cada conjunto de dez é representado pelo quadrado formado por quatro pontos e seis linhas. Para contar os dez seguintes faz-se um novo quadrado. Assim, cada dezena é facilmente reconhecível, sendo igualmente muito mais fácil de contar.



14 | RELAÇÕES

Relação: O conversador civilizado só usa esta palavra em público para descrever uma lista de vários itens.

Cleveland Amory

A maioria das revistas tem uma quantidade quase infinita de artigos e correspondência acerca das relações. Porquê? Resposta: as relações são complicadas, às vezes são interessantes e frequentemente parecem imprevisíveis. Este é exatamente o tipo de situação em que a matemática pode ser-lhe útil.

As relações mais simples entre objetos têm uma propriedade a que chamamos "transitividade" e que simplifica bastante a vida. Ser "mais alto do que" é uma destas relações transitivas. Assim, se Ali é mais alto do que Bob e Bob é mais alto do que Carla, então Ali é necessariamente mais alto do que Carla. A relação é uma propriedade das alturas. Mas nem todas as relações são assim. Ali pode gostar de Bob e Bob pode gostar de Carla, mas isso não significa que Ali goste de Carla. Estas relações "intransitivas" podem criar situações muito invulgares no que diz respeito a decidir o que se deve fazer quando todos os membros do grupo não estão de acordo.

Suponhamos que Ali, Bob e Carla decidem investir dinheiro em conjunto na compra de um carro em segunda mão e consideram três possibilidades distintas: um Audi, um BMW e um Reliant Robin. Não estão todos de acordo em relação ao carro que querem comprar, portanto concluem que a decisão tem de ser tomada democraticamente; têm de fazer uma votação. Assim, cada um deles escreve a sua ordem de preferência das três marcas:

	Primeira escolha	Segunda escolha	Terceira escolha
Ali	Audi	BMW	Reliant Robin
Bob	BMW	Reliant Robin	Audi
Carla	Reliant Robin	Audi	BMW

Inicialmente, a votação parece promissora: o Audi vence o BMW por duas preferências para uma, e o BMW vence o Reliant Robin por duas preferências para uma. No entanto, estranhamente, o Reliant Robin vence o Audi por duas preferências para uma. A preferência, tal como a admiração, é uma relação intransitiva que pode criar estranhos paradoxos se não for usada com cuidado. Pequenas votações para decidir o preferido entre vários candidatos a um emprego ou para decidir quem será o capitão de uma equipa desportiva ou mesmo para escolher o carro a comprar estão carregadas de paradoxos. Os votantes têm de ter cuidado.

Confrontados com este dilema, Ali, Bob e Carla decidiram desistir da compra do carro e investir os seus recursos no aluguer de uma casa juntos. Logo se tornaram necessárias mais decisões. Deveriam decorar a sala de estar? Deveriam tratar do jardim? Deveriam comprar um televisor novo? Como não chegavam a um consenso, decidiram votar "sim" ou "não" para cada uma destas decisões. Eis as respostas obtidas:

	Decorar a casa?	Tratar do jardim?	Comprar TV?
Ali	Sim	Sim	Não
Bob	Não	Sim	Sim
Carla	Sim	Não	Sim
<i>Decisão maioritária</i>	<i>SIM</i>	<i>SIM</i>	<i>SIM</i>

Agora tudo parecia claro. Havia uma decisão maioritária de dois para um em todas as três decisões. Deveriam fazer as três coisas. Mas o dinheiro não era suficiente e o grupo decidiu que precisavam de partilhar a casa com mais duas pessoas, para conseguirem pagar a renda. Depois de fazerem apenas alguns telefonemas, encontraram os seus novos colegas de casa, Dell e Tracy, que rapidamente se mudaram para lá com todos os seus pertences. Claro que o mais justo seria que também fossem incluídos na votação relativa à decoração, à jardinagem e à compra do televisor. Ambos votaram "não" em cada uma das três propostas, ao passo que Ali, Bob e Carla mantiveram as suas decisões anteriores. Agora, gerara-se uma situação muito estranha na casa.

Eis a tabela das decisões, depois de incluídos os "Não" de Dell e de Tracy:

	Decorar a casa?	Tratar do jardim?	Comprar TV?
Ali	Sim	Sim	Não
Bob	Não	Sim	Sim
Carla	Sim	Não	Sim
<i>Decisão maioritária anterior</i>	<i>SIM</i>	<i>SIM</i>	<i>SIM</i>
Dell	Não	Não	Não
Tracy	Não	Não	Não
<i>Decisão global</i>	<i>NÃO</i>	<i>NÃO</i>	<i>NÃO</i>

Verificamos que os votos negativos dos dois novos membros alteraram os resultados em cada uma das decisões. Agora há uma maioria de três para dois quanto às decisões de não decorar a casa, não tratar do jardim e não comprar um televisor. Mas mais surpreendente é o facto de a maioria (Ali, Bob e Carla) ter ficado a perder na votação em duas das três decisões – aquelas em que não votaram “não”. Assim, Ali ficou a perder no que diz respeito à casa e ao jardim, Bob ficou a perder quanto ao jardim e à TV e Carla ficou a perder no que concerne a casa e a TV. Portanto, uma maioria das pessoas (três em cinco) ficou a perder na maioria das decisões (duas em três)!

"Tem de haver vigilância constante para garantir que o jogo organizado não é infiltrado por interesses criminosos que, neste contexto, compreendem operadores sofisticados, inteligentes, altamente organizados, bem informados e esperados com vastos recursos financeiros que lhes permitem contratar os melhores cérebros nas áreas do direito, da contabilidade, da gestão, da hotelaria e do mundo do espetáculo." Tenho algumas dúvidas em relação aos dois últimos; ainda assim, é uma verdade que continua a ser tão atual como naquela altura.

Visconde Falkland a citar o *Report of Rothschild Commission on Gambling* (1979)

Há algum tempo, vi um filme policial na televisão que envolvia um plano para defraudar agentes de apostas, drogando o favorito numa corrida. A intriga centrava-se noutros acontecimentos, como um homicídio, e a base da fraude nas apostas nunca foi explicada. O que se andaria a passar?

Suponhamos que temos uma corrida em que são publicadas as seguintes *odds* para os concorrentes: a_1 para 1, a_2 para 1, a_3 para 1, etc., para qualquer número, N , de competidores numa corrida. Se as *odds* forem de 5 para 4, são expressadas como a_i de 5/4 para 1. Se apostarmos em todos os N concorrentes em proporção às *odds* de modo a apostarmos uma fração $1/(a_i+1)$ do valor total das apostas no corredor com as *odds* de a_i para 1, teremos sempre lucro desde que a soma das *odds*, a que chamaremos Q , satisfaça a desigualdade:

$$Q = 1/(a_1+1) + 1/(a_2+1) + 1/(a_3+1) + \dots + 1/(a_N+1) < 1$$

E se Q for, de facto, inferior a 1, os nossos lucros serão, pelo menos, iguais a:

$$\text{Lucros} = (1/Q - 1) \times \text{o valor total das apostas}$$

Consideremos alguns exemplos. Suponhamos que temos quatro corredores e que as *odds* para cada um deles são de 6 para 1, 7 para 2, 2 para 1 e 8 para 1. Neste caso temos $a_1=6$, $a_2=7$, $a_3=2$ e $a_4=8$ e:

$$Q = 1/7 + 2/9 + 1/3 + 1/9 = 51/63 < 1$$

E assim, ao apostarmos 1/7 no corredor 1, 2/9 no corredor 2, 1/3 no corredor 3 e 1/9 no corredor 4, ganharemos pelo menos 21/51 do dinheiro apostado (para além de recuperarmos o valor apostado).

No entanto, suponhamos que na corrida seguinte as *odds* para os quatro corredores são de 3 para 1, 7 para 1, 3 para 2 e 1 para 1 (ou seja, 50-50). Agora verificamos que temos:

$$Q = 1/4 + 1/8 + 2/5 + 1/2 = 51/40 > 1$$

e não podemos garantir um retorno positivo. Geralmente, podemos ver que se existir um grande grupo de corredores (resultando num número N grande), haverá uma maior probabilidade de Q ser superior a 1. Mas um N grande não garante necessariamente que obtenhamos um $Q > 1$. Basta-lhe escolher cada uma das *odds* segundo a fórmula $a_i = i(i+2)-1$ e obterá $Q = 3/4$ e uma agradável retorno de 30, mesmo que N seja infinito!

Voltemos agora ao filme que vi na televisão. De que forma é que a situação se altera se soubermos antecipadamente que o favorito no nosso exemplo $Q > 1$ não vai participar na corrida por ter sido drogado?

Se usarmos este conhecimento privilegiado de que o concorrente foi drogado, excluiremos o favorito (com as *odds* de 1 para 1) e não apostaremos nele. Assim, estaremos a apostar numa corrida com apenas três cavalos em que Q é igual a:

$$Q = 1/4 + 1/8 + 2/5 = 31/40 < 1$$

E, ao apostarmos 1/4 do nosso dinheiro no corredor 1, 1/8 no corredor 2 e 2/5 no corredor 3 temos garantido um retorno mínimo de $(40/31)-1 = 9/31$ no valor total das nossas apostas para além do valor apostado! Portanto ficamos a ganhar*.

* Sugeriram-me que algumas atividades criminosas de lavagem de dinheiro são efetuadas apostando em todos os concorrentes, mesmo quando $Q > 1$. Neste caso há uma perda, mas de um modo geral é possível prever o valor dessa perda e é encarada como uma espécie de "imposto" sobre a lavagem de dinheiro.

16 | SALTOS EM ALTURA

Há dois tipos de trabalho: um consiste em alterar a posição da matéria no centro da superfície da Terra ou perto dele em relação a outros tipos de matéria; o segundo consiste em mandar alguém fazê-lo. O primeiro tipo é desagradável e mal pago; o segundo é agradável e muito bem pago.

Bertrand Russell

Se estiver a treinar para ser bom em algum desporto, está a trabalhar na atividade da otimização – a fazer tudo o que pode (dentro dos limites da lei) para aperfeiçoar uma característica que o vai tornar melhor e ajudá-lo a minimizar as falhas que reduzem o seu desempenho. Esta é uma das áreas da ciência do desporto que se baseia em conhecimentos tornados possíveis pela aplicação de um pouco de matemática. Existem duas modalidades desportivas em que se tenta lançar o corpo por cima da maior altura possível acima do chão: o salto em altura e o salto à vara. Este tipo de modalidades não é tão simples como pode parecer à primeira vista. Em primeiro lugar, os atletas têm de usar a sua força e energia para lançar o peso do corpo para o ar, desafiando a gravidade. Se pensarmos num praticante de salto em altura como um projétil de massa M lançado na vertical a uma velocidade U , a altura A que pode ser alcançada é determinada pela fórmula $U^2 = 2gA$, em que g corresponde à aceleração produzida pela gravidade. A energia do movimento do saltador no momento do salto é $1/2 MU^2$ e será transformada na energia potencial MgA ganha pelo saltador na altura máxima A . Ao igualar os dois termos, obtemos $U^2 = 2gA$.

A dificuldade está na quantidade A : o que é, exatamente? Não é a altura transposta pelo saltador, mas sim a altura a que o centro de gravidade deste se ergue, o que é uma questão muito subtil na medida em que torna possível que o corpo de um praticante de salto em altura passe por cima da fasquia, apesar de o seu centro de gravidade passar por baixo dela.

Quando um objeto tem uma forma curva, como um L, é possível que o centro de gravidade esteja situado fora do corpo*. É esta possibilidade que permite que um saltador controle a localização do seu centro de gravidade e a trajetória que este deve seguir durante o salto. O objetivo do saltador em altura é fazer o corpo passar por cima da fasquia sem lhe tocar ao mesmo tempo que faz o seu centro de gravidade passar por baixo da fasquia, o mais baixo possível. Desta forma, fará um uso ideal da energia explosiva do lançamento para aumentar a altura a que salta.

O estilo simples de salto em altura que aprendemos na escola, designado por "técnica da tesoura", está longe de ser ideal. Para saltar a fasquia, o seu centro de gravidade, bem como o resto do corpo, tem de passar por cima da fasquia. De facto, o centro de gravidade passa provavelmente a uma altura de 30 centímetros acima da fasquia. É uma forma muito pouco eficiente de fazer um salto em altura.

As técnicas de salto em altura usadas pelos atletas profissionais são muito mais elaboradas. A velha técnica *straddle* envolvia o saltador dar a volta à fasquia sempre com o peito voltado para ela. Esta foi a técnica preferida dos saltadores profissionais até 1968, quando o americano Dick Fosbury surpreendeu toda a gente ao usar uma técnica completamente diferente – que ficou conhecida como o "Fosbury Flop" – que envolvia saltar com as costas voltadas para a fasquia. Esta técnica levou-o a ganhar a medalha de ouro nas Olimpíadas de 1968 na Cidade do México. Este método só se tornou seguro quando foram disponibilizados os colchões insufláveis. A técnica de Fosbury foi muito mais fácil de aprender para os atletas que faziam salto em altura do que a *straddle* e atualmente é usada por todos os bons atletas da modalidade. Permite a um saltador fazer o seu centro de gravidade passar muito abaixo da fasquia, apesar de o corpo se virar e passar por cima dela. Quanto mais flexível uma pessoa for, maior será a sua capacidade de curvar o corpo em volta da fasquia e mais baixo será o seu centro de gravidade. O campeão de salto em altura das Olimpíadas de 2004, Stefan Holm, da Suécia, é bastante baixo em comparação com a maioria dos atletas desta modalidade, mas consegue curvar o corpo

* Uma forma de localizar o centro de gravidade de um objeto é pendurá-lo por uma ponta e prender-lhe um fio com um peso na extremidade, marcando o ponto onde o fio vai cair. Depois, repete-se este passo, pendurando o objeto por outra ponta. Traça-se uma segunda linha onde cai agora o fio. O centro de gravidade é o ponto onde as linhas dos dois fios se cruzam. Se o objeto for quadrado, o centro de gravidade fica situado no seu centro geométrico, mas se tiver a forma de um L ou de um U, de um modo geral o centro de gravidade ficará fora dos limites do corpo.

a um limite absolutamente notável. Quando alcança o ponto mais elevado, o seu corpo forma um U quase perfeito. Consegue passar uma fasquia colocada a 2,37 metros de altura, mas o seu centro de gravidade passa muito baixo da fasquia.

Quando um praticante de salto em altura corre para se elevar no ar, consegue transferir apenas uma pequena fração da sua melhor velocidade de corrida horizontal para o salto vertical. Tem pouco espaço para a corrida de balanço e tem de se virar para ter as costas voltadas para a fasquia quando inicia o salto. O praticante de salto à vara consegue fazer muito melhor. Pode fazer uma corrida longa e em linha reta ao longo do corredor de balanço e, apesar de transportarem uma vara muito comprida, os melhores praticantes de salto à vara do mundo conseguem alcançar velocidades de quase 10 metros por segundo na descolagem. A vara elástica de fibra de vidro permite-lhes transferir a energia do movimento horizontal $1/2 MU^2$ para o movimento vertical de uma forma muito mais eficiente do que a alcançada por um praticante de salto em altura. Os praticantes de salto à vara lançam o corpo para cima na vertical e fazem a espetacular ginástica necessária para enrolar o corpo na forma de um U invertido por cima da fasquia, enquanto mantêm o centro de gravidade o mais baixo possível.

Vejam se conseguimos encontrar uma estimativa aproximada do desempenho esperado deles. Suponhamos que conseguem transformar toda a energia cinética da corrida de $1/2 MU^2$ em energia elástica curvando a vara e depois numa energia vertical potencial de MgA . Conseguirão elevar o centro da sua massa a uma altura de $A = U^2/2g$.

Se o campeão olímpico conseguir alcançar uma velocidade de 9 metros por segundo na descolagem, tendo em conta que a aceleração produzida pela gravidade é $g = 10 \text{ ms}^{-2}$, espera-se que ele consiga erguer o seu centro de gravidade a $A = 4$ metros. Se o saltador tiver começado com o centro de gravidade a 1,5 metros acima do chão e o tiver feito passar a 0,5 metros abaixo da fasquia, espera-se que consiga saltar uma fasquia colocada a uma altura de $1,5 + 4 + 0,5 = 6$ metros. De facto, o campeão americano Tim Mack ganhou a medalha de ouro nos Jogos Olímpicos de Atenas com um salto a 5,95 metros (19'6" em pés e polegadas) e falhou três vezes os 6 metros por muito pouco, sabendo que já tinha conquistado a medalha de ouro, pelo que as nossas estimativas se revelaram bastante exatas.

A periferia é onde o futuro se revela.

J.G. Ballard

Os limites são importantes. E não apenas por manterem os lobos do lado de fora e as ovelhas do lado de dentro. Determinam o nível de interação que pode existir entre uma coisa e outra e o nível de exposição que alguma coisa a nível local pode ter ao mundo exterior.

Tomemos como exemplo um laço fechado de um fio com o comprimento p e pousemo-lo numa mesa. Qual é a área que pode encerrar? Se mudarmos a forma do laço, observaremos que ao alongá-lo ou estreitá-lo podemos diminuir cada vez mais a área encerrada. A área máxima encerrada ocorre quando o laço tem uma forma circular. Nesse caso, sabemos que $p = 2\pi r$ é a área, em que r corresponde ao raio do laço feito de fio. Desta forma, eliminando r , para *qualquer* laço fechado com um perímetro p que encerre uma área A , esperamos que $p^2 \geq 4\pi A$, com a igualdade a surgir apenas no caso de um círculo. Por outras palavras, isto diz-nos que podemos alongar o quanto quisermos o perímetro de uma determinada área encerrada. A melhor forma de o conseguir é torná-la progressivamente mais ondulada.

Se passarmos das linhas que encerram áreas às superfícies que encerram volumes, deparamo-nos com um problema semelhante. Que forma pode maximizar ao máximo o volume contido numa determinada área de superfície? Mais uma vez, o maior volume encerrado é obtido pela esfera, com um volume $V = 4\pi r^3/3$ dentro de uma superfície esférica com uma área $A = 4\pi r^2$. Assim, para qualquer superfície fechada com uma área A , esperamos que o volume encerrado obedeça a $A^3 \geq 36\pi V^2$, com igualdade no caso da esfera. Tal como no caso anterior, observamos que ao tornar

a superfície recortada e ondulada, podemos aumentar progressivamente a área que encerra um determinado volume. Esta é uma estratégia muito bem-sucedida que os sistemas vivos exploram habitualmente.

Há muitas situações em que uma grande área de superfície é importante. Se quiser manter-se fresco, quanto maior a sua área de superfície, melhor. Pelo contrário, se quiser manter-se quente, é desejável uma área de superfície menor – é por isso que as aves e mamíferos recém-nascidos se amontoam numa bola, de modo a minimizar a superfície exposta. Da mesma forma, os membros de uma manada de gado ou de um cardume de peixes que procurem minimizar as oportunidades de serem apanhados pelos predadores oporão por se reunir num grupo com forma circular ou esférica de forma a minimizar a superfície que um predador consegue atacar. No caso de uma árvore, que retira humidade e nutrientes do ar, compensa maximizar a área que está em contacto com a atmosfera, pelo que é vantajoso ter muitos ramos, com muitas folhas onduladas. Se se tratar de um animal que esteja a tentar absorver o máximo possível de oxigénio pelos pulmões, esse objetivo é concretizado de forma mais eficaz através da maximização dos canais que podem estar contidos no volume dos pulmões, de modo a maximizar a superfície que interage com as moléculas de oxigénio. Se pretender simplesmente secar-se quando sai do duche, uma toalha com uma grande superfície será mais adequada. Por este motivo, as toalhas de banho tendem a ter uma superfície coberta de pequenas franjas: a superfície por unidade de volume é muito maior do que se fossem lisas. Esta luta para maximizar a superfície que encerra um volume é algo que vemos constantemente no mundo natural. É a razão pela qual as "fractais" aparecem constantemente como uma solução evolutiva para os problemas da vida. Facultam a forma sistemática mais simples de obter uma superfície superior à que seria esperada para o volume em questão.

É preferível manter-se num único grupo ou separar o grupo em dois ou mais grupos menores? Este problema foi enfrentado pelos comboios de navios que tentavam evitar ser detetados pelos submarinos inimigos durante a Segunda Guerra Mundial.

É preferível permanecer num comboio grande do que dividir-se em vários mais pequenos. Suponhamos que esse comboio cobre uma área A e que os navios estão o mais juntos possível, de modo que se dividíssemos o comboio em dois mais pequenos com uma área de $A/2$, o espaço entre

os navios seria o mesmo. O comboio único tem um perímetro igual a $p=2\pi\sqrt{(A/\pi)}$, mas o perímetro total dos dois comboios mais pequenos é igual a $p \times \sqrt{2}$, que é maior. Assim, a distância total do perímetro que tem de ser patrulhado pelos contratorpedeiros para proteger os dois comboios mais pequenos dos ataques dos submarinos é superior àquela que teria de ser patrulhada se se tivesse mantido um comboio único. Além disso, quando o submarino procura comboios para atacar, a sua probabilidade de os ver é proporcional ao diâmetro deles, pois é essa a imagem que se vê pelo periscópio. O diâmetro de um único comboio circular com uma área A é apenas de $2\sqrt{(A/\pi)}$, ao passo que a soma dos diâmetros dos dois comboios de área $A/2$ que se sobrepõem no campo de visão é superior por um fator de $\sqrt{2}$, e assim o comboio dividido tem 41% mais de probabilidades de ser detetado pelo submarino atacante do que o comboio único.

18 | IVA ATÉ À ETERNIDADE

As únicas coisas garantidas neste mundo são a morte e os impostos.

Benjamin Franklin

Se o leitor viver no Reino Unido saberá que existe um imposto somado à venda de muitos produtos e que é designado por “Value Added Tax” ou VAT. Na Europa continental é frequentemente designado por IVA. No Reino Unido corresponde a 17,5% somados ao preço de uma variedade de produtos e serviços e é a maior fonte de receita fiscal do Estado. Se partirmos do pressuposto de que a taxa de IVA de 17,5% foi criada para permitir um cálculo fácil através de um exercício de aritmética mental, qual podemos esperar que seja o próximo aumento da taxa de IVA? E qual será a taxa de IVA no futuro infinito?

A taxa de IVA atual parece estranha – porquê escolher uma taxa de 17,5%? No entanto, uma pessoa que tenha de preparar os relatórios trimestrais de IVA para uma pequena empresa reconhecerá rapidamente as vantagens deste estranho número. Permite um cálculo mental muito simples, porque $17,5\% = 10\% + 5\% + 2,5\%$, pelo que conseguimos obter facilmente o valor de 10% (basta avançar uma casa decimal para a esquerda), e depois basta encontrar metade desses 10% e depois metade dos 5% restantes e somar os três números. Assim, o IVA de uma compra de 80£, por exemplo, é $8£ + 4£ + 2£ = 14£$.

Se se mantiver a mesma estrutura conveniente de divisão por dois para a aritmética mental, o próximo aumento da taxa de IVA será de metade de 2,5%, ou seja, 1,25%, o que produzirá uma taxa total de 18,75% e o novo valor de IVA correspondente a uma compra de 80£ será $8£ + 4£ + 2£ + 1£ = 15£$.

Teremos uma taxa de valor acrescentado equivalente a $10\% + 5\% + 2,5\% + 1,25\%$. Aos olhos de um matemático, isto parece o início de uma

série interminável em que a próxima parcela da soma é sempre metade da anterior. A taxa de IVA atual é simplesmente:

$$10\% \times (1 + \frac{1}{2} + \frac{1}{4})$$

Se prolongarmos esta série até ao infinito, poderemos prever que a taxa de IVA no futuro infinito será igual a

$$10\% \times (1 + \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \frac{1}{16} + \frac{1}{32} + \dots) = 10\% \times 2$$

que, conforme poderemos constatar na página 252, mostra que a mesma série, excetuando a primeira parcela, somará 1, pelo que a soma de uma série infinita entre parênteses será 2. A taxa de IVA sujeita a aumentos ao longo de um período infinito será de 20%!

19 | VIVER NUMA SIMULAÇÃO

Nada é real.

The Beatles, "Strawberry Fields Forever"

Será a cosmologia um caminho traiçoeiro na direção da ficção científica? As novas observações via satélite das radiações fósseis do Big Bang vieram confirmar a teoria preferida de muitos cientistas acerca da evolução do Universo. No entanto, isto pode não ser exatamente uma boa notícia.

O modelo preferido contém muitas "coincidências" aparentes que permitem que o Universo suporte a complexidade e a vida. Se considerássemos considerar o "multiverso" de todos os universos possíveis, concluiríamos que o nosso é especial em vários aspetos. A física quântica moderna até explica de que formas estes universos possíveis que constituem o "multiverso" podem existir.

Se levarmos a sério a sugestão de que todos os universos possíveis podem existir (ou existem realmente), também teremos de encarar uma consequência bastante estranha. Neste conjunto infinito de universos existirão civilizações tecnologicamente muito mais avançadas do que a nossa, que terão a capacidade de simular universos. Em vez de simularem apenas as condições climatéricas destes ou a formação das galáxias – como nós fazemos –, estas civilizações poderão ir ainda mais longe e estudar a formação das estrelas e dos sistemas planetários. Então, tendo adicionado as regras da bioquímica às suas simulações astronómicas, poderiam assistir à evolução da vida e da consciência nas suas simulações por computador (aceleradas para a evolução ocorrer na escala de tempo que considerassem conveniente). Tal como nós observamos os ciclos da vida das moscas da fruta, eles poderiam seguir a evolução da vida humana, ver as nossas civilizações a crescerem e a comunicarem entre si, e até vê-las a

discutir a hipótese de ter havido um Grande Programador no Céu que criou o seu Universo e que pode intervir a seu bel-prazer, desafiando as leis da natureza que observam habitualmente.

Nestes universos, podem emergir entidades autoconscientes que comunicam entre si. Depois de alcançada essa capacidade, os universos fabricados proliferarão e em pouco tempo serão mais do que os universos reais. Os simuladores determinam as leis que governam estes mundos artificiais; podem criar aperfeiçoamentos que ajudem à evolução das formas de vida que mais lhes agradem. E podemos acabar por nos deparar com um cenário em que, estatisticamente, é mais provável que estejamos a viver numa realidade simulada do que no mundo real, uma vez que os mundos simulados são muito mais do que os reais.

O físico Paul Davies sugeriu recentemente que esta elevada probabilidade de estarmos a viver numa realidade simulada é a prova pelo absurdo da falsidade da existência de um “multiverso”. Mas, defrontados com este cenário, temos forma de apurar a verdade? Talvez, se analisarmos a situação atentamente.

Para começar, os simuladores terão tentado evitar a complexidade de usar um conjunto consistente de leis da natureza nos seus mundos, uma vez que conseguem simplesmente colar neles os efeitos “realistas”. Quando a Disney faz um filme em que é retratada a luz refletida na superfície de um lago, não usa as leis da eletrodinâmica quântica e da ótica para calcular a dispersão da luz. Isso exigiria uma quantidade impressionante de detalhes e de capacidade informática. Em vez disso, a simulação do efeito de dispersão da luz é substituída por regras plausíveis que são muito mais abreviadas do que as reais, mas que produzem um resultado com um efeito realista – desde que ninguém observe com muita atenção. Haveria um imperativo prático e económico que permitiria fazer as realidades simuladas permanecerem assim se fossem destinadas apenas a fins de entretenimento. Mas estas limitações da complexidade da programação da simulação iriam, supostamente, causar ocasionalmente problemas aparentes – que talvez até fossem visíveis a partir do interior.

Mesmo que os simuladores fossem bastante escrupulosos no que diz respeito à simulação das leis da natureza, haveria um limite para o que poderiam fazer. Supondo que os simuladores, ou pelo menos as primeiras gerações de simuladores, tinham um conhecimento muito avançado das

leis da natureza, é provável que tivessem ainda assim um conhecimento incompleto delas (algumas filosofias da ciência argumentam que este é, forçosamente, o caso). Podem saber muito acerca da física e da programação necessárias para simular um universo, mas teriam lacunas ou, pior, erros no seu conhecimento das leis da natureza. Estas lacunas e erros seriam, obviamente, subtis e não seriam de forma alguma óbvios aos nossos olhos, caso contrário esta civilização “avançada” não seria muito avançada. Estas lacunas não impedem a criação e o bom funcionamento das simulações durante longos períodos, mas gradualmente as pequenas falhas começam a acumular-se.

Os efeitos destas falhas acabariam por formar um efeito “bola de neve” e as realidades simuladas bloqueariam. A única solução seria os criadores intervirem e resolverem os problemas um a um, à medida que surgissem. Esta solução é muito familiar para qualquer pessoa que tenha um computador doméstico e que receba atualizações regulares destinadas a protegê-lo de novos ataques de vírus ou a reparar falhas que não foram previstas pelos programadores. Os criadores de uma simulação poderiam oferecer este tipo de proteção temporária, atualizando as leis da natureza em funcionamento para incluir extras que só descobriram depois de ter sido iniciada a simulação.

Neste tipo de situações, surgirão inevitavelmente contradições lógicas e as leis das simulações parecerão falhar ocasionalmente. Os habitantes da simulação – especialmente os cientistas simulados – ficarão ocasionalmente confusos com as observações que fazem. Os astrónomos simulados poderão, por exemplo, fazer observações que demonstram que as ditas constantes da natureza estão a mudar lentamente.

É provável que até surjam falhas súbitas nas leis que governam estas realidades simuladas. Isto porque os simuladores utilizariam provavelmente uma técnica que foi considerada eficaz em todas as outras simulações de sistemas complexos: o uso de códigos de correção de erros que poriam tudo a funcionar novamente.

Consideremos, por exemplo, o nosso código genético. Se só dependesse de si próprio, nós não duraríamos muito. Os erros acumular-se-iam e seguir-se-iam rapidamente a morte e as mutações. Somos protegidos de uma situação deste tipo pela existência de um mecanismo de correção de erros que identifica e corrige os erros na codificação genética. Muitos dos nossos

sistemas informáticos complexos possuem o mesmo tipo de sistema imunitário interno que impede a acumulação de erros.

Se os simuladores usassem códigos informáticos para a correção de erros de modo a combater a falibilidade das suas simulações como um todo (bem como simulá-las numa escala menor no nosso código genético), ocasionalmente ocorreria uma correção no estado ou nas leis que governam a simulação. Ocorreriam mudanças misteriosas que pareceriam contrariar as próprias leis da natureza que os cientistas simulados estavam habituados a observar e a prever.

Parece, portanto, tentador concluir que, se vivermos numa realidade simulada, podemos esperar "falhas" ocasionais ou resultados experimentais que não conseguimos repetir, ou mesmo lentos desvios das supostas constantes e leis da natureza que não conseguimos explicar.

Um político precisa de ser capaz de prever o que vai acontecer amanhã, na próxima semana, no próximo mês e no próximo ano. E de ter a capacidade de depois explicar por que é que as suas previsões não se concretizaram.

Winston Churchill

Uma das palavras da moda nas ciências que estudam as coisas complicadas é “emergência”. Quando se constrói uma situação complexa passo a passo, parece que se podem alcançar fronteiras da complexidade que anunciam o aparecimento de novas estruturas e novos tipos de comportamento que não estavam presentes nas etapas iniciais dessa complexidade. A Internet ou a bolsa de valores ou a consciência humana podem parecer fenómenos deste tipo. Apresentam um comportamento coletivo que é mais do que a soma das suas partes. Se as reduzirmos aos seus componentes elementares, a essência do comportamento complexo desaparece. Este tipo de fenómenos também é comum na física. Uma propriedade coletiva de um líquido, como a viscosidade, que descreve a sua resistência ao fluxo, emerge quando um grande número de moléculas se combina. É real e, no entanto, não conseguirá encontrar a menor viscosidade numa molécula de hidrogénio e oxigénio na sua chávena de chá.

A emergência é, em si, um assunto complexo e, ocasionalmente, controverso. Os filósofos e os cientistas tentam defini-la e distinguir diferentes tipos de emergência, e há ainda alguns que contestam a sua existência. Um dos problemas é que os exemplos científicos mais interessantes, como a consciência ou a “vida”, não são compreendidos e, assim, há uma camada adicional de incerteza somada aos casos usados como exemplos. Aqui, a matemática pode ajudar. Dá origem a muitas interessantes estruturas emergentes que são bem definidas e que sugerem formas de criar famílias inteiras de novos exemplos.

Consideremos conjuntos finitos de números positivos como $[1, 2, 3, 6, 7, 9]$. Por maior que seja o conjunto, não possuirá as propriedades que “emergem” quando o conjunto se torna infinito. Como Georg Cantor demonstrou claramente no século XIX, conjuntos infinitos de números possuem propriedades que não são partilhadas por nenhum subconjunto finito desses números, por maior que seja. A infinidade *não é* apenas um número muito grande. Pode somar-lhe 1 e permanece o mesmo; subtraia-lhe a infinidade e permanece o mesmo. O todo não é apenas maior do que as partes, possui também características “emergentes” qualitativamente diferentes das de qualquer uma das suas partes.



Banda de Möbius

É possível encontrar muitos outros exemplos na topologia, em que a estrutura global de um objeto pode ser surpreendentemente diferente da sua estrutura local. O exemplo mais familiar é a banda de Möbius. Pode ser construída pegando numa tira retangular de papel e colando as pontas depois de o torcer uma vez sobre si próprio. Também é possível construir essa tira de papel colando pequenos retângulos de papel, como uma manta de retalhos. Neste caso, a banda de Möbius parecerá um tipo de estrutura emergente. Todos os retângulos que foram unidos para formar essa banda têm duas faces, mas quando as pontas são torcidas e depois coladas, a banda de Möbius resultante passa a ter apenas uma face. Mais uma vez, o todo tem uma propriedade que não é partilhada pelas suas partes.

21 | COMO EMPURRAR UM CARRO

Há apenas dois tipos de peões no trânsito intenso dos dias de hoje – os rápidos e os mortos.

Lord Dewar

Antigamente contava-se a seguinte piada: "Por que é que o carro Lada tem um aquecedor voltado para o para-brisas traseiro?" Resposta: "Para não ficarmos com as mãos frias quando o empurrarmos." Mas o ato de empurrar um carro apresenta-nos um problema interessante. Suponhamos que tem de empurrar o seu carro para dentro da garagem e fazê-lo parar antes de ele bater na parede do fundo. Como é que deve empurrar de forma a fazer o carro entrar na garagem e puxar para parar o mais rapidamente possível?

A resposta é que deve acelerar o carro empurrando com toda a sua força durante metade da distância desejada e depois fazê-lo desacelerar, puxando com toda a sua força durante a outra metade. O carro começará por estar parado e chegará ao local desejado parado, e conseguirá fazê-lo no menor espaço de tempo possível².

Este tipo de problemas é um exemplo de uma área da matemática chamada "teoria do controlo". Habitualmente, tentamos regular ou guiar algum tipo de movimento através da aplicação de uma força. A solução para o problema do estacionamento do carro é um exemplo daquilo a que se chama o controlo do tipo *bang-bang*. Há apenas duas respostas: empurrar e depois puxar. Os termóstatos domésticos têm frequentemente um funcionamento semelhante. Quando a temperatura sobe demasiado, ligam o sistema de refrigeração e quando a temperatura desce demasiado, ligam o aquecimento. Ao fim de um longo período, a temperatura alterna entre os dois extremos que programou. Esta não é sempre a melhor forma de controlar uma situação. Suponhamos que quer con-

trolar o seu carro numa estrada usando o volante. Um condutor-robô programado com uma abordagem ao controlo do tipo *bang-bang* deixaria o carro ir para a faixa da esquerda e depois corrigiria a trajetória para a direita até passar o traço que divide a estrada e entrar na faixa da direita, e continuaria a repetir este sistema, para um lado e para o outro. Se usasse a estratégia *bang-bang* na condução, dentro de pouco tempo, o condutor seria parado pela polícia e convidado a soprar para dentro de um tubo de plástico antes de ser preso. Uma abordagem melhor seria aplicar correções proporcionais ao grau de desvio da posição central. É assim que funcionam os baloiços. Se o empurrar levemente para fora da posição vertical, ele recuará mais lentamente do que se for afastado dessa posição com um empurrão brusco.

Outra aplicação interessante da teoria do controlo foi o estudo das corridas de fundo e de semi-fundo — que supostamente também seria aplicável às corridas de cavalos, para além de às corridas humanas. Tendo em conta que existe uma quantidade limitada de oxigénio disponível nos músculos de um atleta e um limite para a quantidade de oxigénio que pode ser restabelecido pela respiração, qual é a melhor forma de correr de modo a minimizar o tempo necessário para percorrer uma determinada distância? Uma solução do tipo *bang-bang* baseada na teoria do controlo especifica que para distâncias superiores a 300 metros (que sabemos que é onde começa o exercício anaeróbico e o débito de oxigénio), diz-nos que se deveria aplicar a máxima aceleração durante um breve período e depois correr a uma velocidade constante antes de desacelerar no fim da corrida durante o mesmo período que durou a aceleração inicial. Claro que se por um lado isto pode ser indicativo de como melhor correr um contrarrelógio, talvez não seja a melhor forma de ganhar uma corrida em que está a competir com outros atletas. Se tiver um bom *sprint* final ou se tiver treinado para lidar com mudanças muito drásticas de ritmo, poderá ter uma vantagem em relação aos outros ao adotar uma tática diferente. Seria uma grande demonstração de coragem escolher a solução ideal quando compete com outros atletas. Posicionar-se atrás do atleta que opta pela estratégia ideal, mantendo-se abrigado do vento e ganhando propulsão adicional antes de acelerar para a vitória na reta final é um bom plano alternativo.

Acentue o positivo. Elimine o negativo. Agarre-se ao afirmativo. Fuja do intermediário.

Johnny Mercer e Harold Arlen, "Ac-cent-tchu-ate the Positive"

No início deste ano tive uma experiência estranha enquanto estava hospedado num hotel em Liverpool durante um nevão inesperado. O hotel tinha sido construído no século XIX, no apogeu comercial da cidade, e reconstruído ao estilo "boutique". De manhã tinha percorrido um caminho tortuoso no meio da neve, vindo de Manchester, num comboio em marcha lenta que acabou por ser parado por um longo período depois de o maquinista ter anunciado que o cabo de sinalização ao longo da extensão de caminho de ferro que tínhamos pela frente fora roubado durante a noite – mais uma indicação do aumento do valor do cobre, observei. Por fim, depois de vários telefonemas, o comboio avançou lentamente, passando por todos os sinais que tinham ficado vermelhos por predefinição, e chegou em segurança à estação de Lime Street.

O meu quarto no hotel estava frio e a temperatura no exterior da abóbada coberta de neve era inferior a zero. O aquecimento estava instalado debaixo do chão, tendo um tempo de atuação bastante lento, e tive dificuldade em perceber se estava a reagir à alteração das definições do termóstato. Apesar das garantias dos funcionários de que a temperatura não tardaria a aumentar, o espaço parecia estar a ficar gradualmente mais frio e acabaram por me trazer um aquecedor com ventilador para ajudar. Alguém na receção sugeriu que o fizéssemos porque o aquecimento era novo e era importante não aumentar demasiado a temperatura.

Muito mais tarde no mesmo dia, a gerência do hotel chamou um técnico, por estar tão preocupada com o mito de que não se pode ligar no máximo um equipamento novo e também por estar confusa por o aque-

cimento parecer estar a funcionar perfeitamente nos corredores do lado de fora do quarto – de tal forma que eu acabei por deixar a porta do quarto aberta. Felizmente, o painel principal do aquecimento para os quartos ficava mesmo em frente à porta do meu quarto, pelo que pudemos ver juntos o que este revelava antes de o técnico ter ido investigar a temperatura do quarto ao lado, cujo hóspede estaria ausente durante todo o dia. O quarto ao lado estava muito quente.

Subitamente, o técnico percebeu o que estava a causar o nosso problema. O aquecimento do quarto contíguo tinha sido ligado ao termóstato do meu quarto e o meu ao daquele. O resultado foi um belíssimo exemplo daquilo a que os engenheiros chamam “instabilidade térmica”. Quando o meu vizinho achava que tinha o quarto demasiado quente, baixava o termóstato. Consequentemente, o meu quarto ficava mais frio, o que me levava a aumentar o meu termóstato, tornando o quarto dele ainda mais quente, pelo que ele baixava mais a temperatura, deixando o meu ainda mais frio e levando-me a aumentar ainda mais o termóstato... Felizmente, o outro hóspede acabou por desistir e saiu.

Este tipo de instabilidade mantém-se com base no interesse próprio de dois intervenientes. O mesmo tipo de situação pode originar problemas ambientais muito mais graves. Se usarmos demasiadas ventoinhas e aparelhos de ar condicionado para nos mantermos frescos, aumentaremos os níveis de dióxido de carbono na atmosfera, o que levará a uma maior retenção do calor do Sol em torno da Terra, aumentando assim a nossa necessidade de baixar ainda mais a temperatura. No entanto, este problema não pode ser resolvido com uma simples retificação das ligações do sistema.

23 | O PASSO DO BÊBEDO

Mostra-me o caminho de casa.
Estou cansado e quero ir dormir.
Bebi um pouco há cerca de uma hora
e subiu-me à cabeça.

Irving King

Um dos testes de sobriedade usados pelas forças policiais em todo o mundo é a avaliação da capacidade de andar em linha reta. Regra geral, é uma tarefa simples para qualquer pessoa em condições físicas normais. E se souber qual é o comprimento médio dos seus passos, saberá exatamente que distância terá percorrido ao fim de uma determinada quantidade de passos. Se o seu passo tiver um metro de comprimento, ao fim de P passos terá avançado P metros desde o ponto de partida. Mas imagine que por um motivo qualquer não consegue caminhar a direito. Aliás, suponhamos que não faz a mínima ideia do que está a fazer. Antes de dar o próximo passo, escolha aleatoriamente a direção, para ter probabilidades iguais de escolher qualquer uma, e avance um metro nessa direção. Em seguida, escolha aleatoriamente outra direção e dê mais um passo. Continue a usar este método para escolher a direção do seu próximo passo e descobrirá que estes traçam um percurso ondulante e bastante imprevisível que foi designado por o Passo do Bêbedo.

Uma questão interessante em relação ao Passo do Bêbedo é a distância em que resultaria se fosse medido em linha reta desde o ponto de partida e depois de o bêbedo ter dado P passos. Sabemos que uma pessoa sóbria precisa de dar P passos de um metro para percorrer em linha reta uma distância P , mas normalmente um bêbedo precisaria de dar P^2 passos³ para alcançar a mesma distância a partir do ponto inicial. Então, 100 passos sóbrios levá-lo-ão a percorrer uma distância de 100 metros,

mas o bêbedo precisará de dar 10 000 passos para percorrer a mesma distância.

Felizmente, esta análise numérica tem mais implicações do que sugere o exemplo dos passos de um caminhante bêbedo. A imagem de uma sucessão aleatória de passos é um modelo excelente do que acontece quando as moléculas se espalham de um lugar para outro por terem uma temperatura superior às que as rodeiam, avançando assim mais depressa do que a média. Colidem com as outras moléculas, umas a seguir às outras, de forma aleatória e espalham-se a partir do ponto inicial como o bêbedo cambaleante, precisando de cerca de N^2 colisões para percorrerem uma distância igual a N das distâncias que as moléculas percorrem entre colisões. É por isso que demoramos uma quantidade significativa de tempo a sentir os efeitos de um aquecedor acabado de ligar numa divisão. As moléculas com mais energia ("mais quentes") estão a deambular como bêbedos pelo espaço, ao passo que as ondas sonoras que emanam da canalização do sistema de aquecimento se propagam em linha reta à velocidade do som.

Blackadder: É o mesmo plano que usámos da última vez e das dezassete vezes anteriores.

Melchett: E-E-Exato! É por isso que é um plano brilhante! Vamos apanhar o guarda completamente desprevenido! A última coisa que eles esperam é que façamos exatamente o mesmo que nas dezoito vezes anteriores!

Richard Curtis e Ben Elton, in *Blackadder*

Uma das ideias falsas mais interessantes em estatística é a de como se processa a distribuição aleatória. Suponhamos que tinha de decidir se uma determinada sequência de acontecimentos era ou não aleatória. Poderia esperar descobrir se ela era aleatória encontrando um padrão ou qualquer outra característica previsível. Vamos tentar inventar algumas listas de “faces” (F) e “versos” (V) que se supõe serem os resultados de uma sequência de lançamentos de uma moeda, de modo a que ninguém consiga distingui-los de lançamentos reais. Aqui estão três possíveis sequências falsas de 32 lançamentos:

VFFVFVFVFVFVFVFVFVVVFVFVFVFVFVF
VFFVFVFVFVFVFVFVFVFVFVFVFVFVFVF
VFVFVFVFVFVFVFVFVFVFVFVFVFVFVF

Parece-lhe certo? Considerá-las-ia sequências aleatórias reais de “faces” e “versos” obtidas com o lançamento de uma moeda, ou parecem apenas uma má falsificação? Em comparação, aqui ficam mais três sequências de “faces” e “versos” para escolher:

* Neste texto optou-se pelas designações “face” e “verso” em vez de “cara” e “coroa”, para se poderem obter iniciais diferentes para as sequências. (N. T.)

VFFFV VVVVFV VFFFV VFVFV VFFF
 VVVVF VVFVF VFFFV VVVVF VVVVF VVVVF
 VVFVF VFVVV VVVVF VVFVF VVVVF VVVVF

Se perguntasse a qualquer pessoa se este segundo grupo de seqüências eram seqüências aleatórias verdadeiras, essa pessoa responderia provavelmente que não. As primeiras três seqüências são muito mais semelhantes àquilo que consideramos aleatório. Há uma alternância muito maior entre faces e versos, e não há longas seqüências de faces e de versos, como as que se veem no segundo grupo. Se o leitor se limitasse a usar o teclado do computador para digitar uma seqüência aleatória de Fs e Vs, teria tendência para alternar muito mais e para evitar seqüências longas da mesma letra, pois, caso contrário, “pareceria” que estava a acrescentar deliberadamente um padrão correlacionado.

Surpreendentemente, o segundo conjunto de seqüências é que corresponde aos resultados de um processo verdadeiramente aleatório. Os primeiros três, com os seus padrões em *staccato* e ausência de longas seqüências de faces ou versos, são os falsos. Não pensamos que uma seqüência aleatória pode ter sucessões muito longas de faces ou versos, mas a sua presença é justamente o que garante a autenticidade de uma seqüência aleatória. O processo de lançamento de moedas não tem memória. A probabilidade de obter face ou verso é sempre $\frac{1}{2}$, independentemente do resultado do último lançamento. Todos os lançamentos são acontecimentos independentes. Desta forma, a probabilidade de uma sucessão de s faces ou s versos surgir em seqüência é simplesmente determinada pela multiplicação $\frac{1}{2} \times \frac{1}{2} \times \frac{1}{2} \times \frac{1}{2} \times \dots \times \frac{1}{2}$, s vezes. Ou seja, $\frac{1}{2}^s$. Mas se lançarmos a nossa moeda N vezes de modo a termos N pontos de partida possíveis para uma seqüência de faces ou versos, a nossa probabilidade de obtermos uma seqüência de comprimento s aumenta para $N \times \frac{1}{2}^s$. Uma seqüência de comprimento s será provável quando $N \times \frac{1}{2}^s$ for aproximadamente igual a 1 – ou seja, quando $N = 2^s$. O significado disto é simples. Se analisarmos uma lista de N lançamentos de uma moeda, podemos esperar encontrar seqüências com um comprimento s em que $N = 2^s$.

Todas as seis seqüências aqui apresentadas tinham um comprimento $N = 32 = 2^5$, pelo que, se tiverem sido geradas aleatoriamente, esperamos que haja uma forte probabilidade de conterem uma seqüência de 5 faces

ou versos e conterão quase certamente sequências de comprimento 4. Por exemplo, com 32 lançamentos, há 28 pontos de partida que nos permitem obter uma sucessão de 5 faces ou versos e, em média, duas sequências de cada são uma forte probabilidade. Desta forma, quando o número de lançamentos aumenta, podemos esquecer a diferença entre o número de lançamentos e o número de pontos de partida e usar a prática regra $N = 2^s$. A ausência destas sequências de faces ou versos é que nos deve fazer desconfiar da autenticidade das primeiras três sequências e fazer-nos confiar na autenticidade do segundo grupo de três. A lição que aprendemos aqui é que as nossas intuições acerca da aleatoriedade são influenciadas pela nossa ideia de que deveria ter um aspeto muito mais ordenado do que na realidade tem. Esta inclinação é manifestada pela nossa expectativa de que os extremos, como as longas sequências do mesmo resultado, não devem acontecer — temos a ideia de que essas sequências são ordenadas por estarem a produzir sempre o mesmo resultado.

Também será interessante ter em mente estes resultados quando consideramos longas sequências de resultados desportivos entre equipas que se defrontam regularmente, *Army vs. Navy*, *Oxford vs. Cambridge*, *Arsenal vs. Spurs*, *AC Milan vs. Inter*, *Lancashire vs. Yorkshire*. Existem frequentemente séries de vitórias em que uma equipa ganha durante muitos anos de seguida, embora nestes casos o efeito não costume ser aleatório: a equipa é formada pelos mesmos jogadores ao longo de vários anos e depois estes deixam de jogar ou saem da equipa e cria-se uma nova equipa.

* É fácil generalizar o resultado de sequências aleatórias em que o número de resultados igualmente prováveis é superior a dois (como os F e V do nosso exemplo). No caso do lançamento de um dado não viciado, a probabilidade de qualquer resultado isolado é de $1/6$ e para obtermos uma sequência do mesmo resultado s vezes, deveríamos esperar ter de lançar cerca de 6^s vezes: mesmo no caso de valores mais baixos para s , é um número muito grande.

Podem fazer-se estatísticas para provar tudo – até a verdade.

Noël Moynihan

As médias são uma coisa engraçada. Pergunte ao estatístico que se afogou num lago com uma profundidade média de 3 centímetros. No entanto, são tão familiares e aparentemente tão claras que confiamos completamente nelas. Mas será certo fazê-lo? Imaginemos dois jogadores de críquete. Chamar-lhes-emos, hipoteticamente, Flintoff e Warne. Estão a jogar uma partida decisiva que determinará o resultado do campeonato. Os patrocinadores ofereceram grandes prémios em dinheiro para os melhores lançadores e batedores do jogo. Flintoff e Warne não estão preocupados com o seu desempenho como batedores – exceto na medida em que querem certificar-se de que não há um resultado melhor por parte do adversário – e estão a dar tudo por tudo para ganharem o grande prémio pelo melhor lançamento.

No primeiro *innings*, Flintoff ganha alguns dos primeiros *wickets*, mas é substituído depois de um longo período de lançamentos muito contidos e termina com uma pontuação de 3 em 17 *runs*, o que dá uma média de 5,67. Então, o lado de Flintoff ocupa o lugar do batedor, e Warne está em muito boa forma, obtendo uma sucessão de *wickets* que lhe garante a pontuação final de 7 em 40, uma média de 5,71 *runs* por cada *wicket*. Mas Flintoff tem a melhor média de lançamentos (a mais baixa), no primeiro *innings*, 5,67 contra 5,71.

No segundo *innings*, Flintoff mostra-se pouco contido, mas rapidamente prova que é imbatível para os batedores inferiores da equipa adversária, e obtém 7 *wickets* em 110 *runs*, uma média de 15,71 para o segundo *innings*. É a vez de Warne lançar a bola para a equipa de Flintoff no último

innings da partida. Não é tão bem sucedido como no primeiro *innings*, mas obtém, ainda assim, 3 *wickets* em 48 *runs*, alcançando uma média de 16,0. Desta forma, Flintoff obtém também a melhor média de lançamentos no segundo *innings*, desta vez alcançando uma vantagem de 15,71 para 16,0.

Lançador	Pontuação do primeiro <i>innings</i>	Média do primeiro <i>innings</i>	Pontuação do segundo <i>innings</i>	Média do segundo <i>innings</i>	Pontuações combinadas	Média combinada
Flintoff	3 em 17	5,67	7 em 110	15,71	10 em 127	12,7
Warne	7 em 40	5,71	3 em 48	16	10 em 88	8,8

Quem é que deveria ganhar o prémio de melhor lançador do jogo? Flintoff teve a melhor média no primeiro *innings* e a melhor média no segundo *innings*. Claramente é ele o vencedor. Mas o patrocinador adotou outra perspectiva e decidiu considerar os números globais da partida. Nos dois *innings*, Flintoff obteve 10 *wickets* em 127 *runs*, com uma média de 12,7 *runs* por *wicket*. Warne, por outro lado, obteve 10 *wickets* em 88 *runs* e teve uma média de 8,8. Warne tem claramente a melhor média e ganha o prémio de melhor lançador, apesar de Flintoff ter tido uma média superior no primeiro e no segundo *innings*!

Podemos lembrar-nos de vários tipos de exemplos semelhantes. Imaginem duas escolas que são classificadas pelos resultados médios dos seus alunos nos exames GCSE*. É possível que uma escola tivesse uma classificação melhor do que a outra em cada uma das disciplinas se fossem analisadas separadamente, mas tivesse uma classificação inferior à da outra escola se fossem considerados os resultados somados de todas as disciplinas. A primeira escola estaria correta se informasse os pais dos alunos de que era superior à outra escola em todas as disciplinas, mas a outra escola podia (também legitimamente) dizer aos pais que os seus alunos tinham, em média, notas mais altas do que os da primeira escola.

As médias são, de facto, uma coisa muito engraçada. *Caveat emptor*: "O cliente que se cuide."

* GCSE (General Certificate of Secondary Education) é um exame feito ao fim de 9, 10 ou 11 anos de escolaridade para dar acesso ao ensino secundário, que corresponde aos anos 12 e 13 no sistema educativo britânico. (N. R.)

Qualquer universo suficientemente simples para ser compreendido, é demasiado simples para produzir uma mente que possa compreendê-lo.

Princípio da Incerteza de Barrow

Se quiser ganhar uma aposta a qualquer adolescente com excesso de confiança, desafie-o a dobrar uma folha de papel A4 mais de 7 vezes. Ele nunca conseguirá fazê-lo. Duplicar e dividir por dois são processos que avançam muito mais depressa do que imaginamos. Suponhamos que pegamos na nossa folha de papel A4 e a dobramos ao meio, e depois novamente ao meio, usando um raio laser para não enfrentarmos problemas na dobragem. Ao fim de apenas 30 dobras, temos uma superfície com 10^{-8} centímetros, aproximadamente o tamanho de um átomo de hidrogénio. Se continuarmos a dobrá-la, ao fim de apenas 47 dobras ficamos com 10^{-13} centímetros, o diâmetro de um único próton que forma o núcleo de um átomo de hidrogénio. Se continuasse a dividir a folha, ao fim de 114 cortes ela alcançaria uma dimensão notável, cerca de 10^{-33} de um centímetro, inimaginável nas nossas unidades métricas antropocêntricas, mas que não é assim tão difícil de conceber quando pensamos em cortar uma folha de papel ao meio apenas 114 vezes, que é certamente muito, mas está longe de ser inimaginável. O que esta escala tem de tão notável é o facto de marcar para os físicos a escala a que as noções de espaço e tempo começam a dissolver-se. Não há teorias da física, nem descrições do espaço, do tempo e da matéria que consigam explicar-nos o que acontece àquele fragmento de papel quando é dividido ao meio apenas 114 vezes. É provável que o espaço, tal como o conhecemos, deixe de existir e seja substituído por alguma forma de “espuma” quântica caótica, em que a gravidade desempenha um papel totalmente novo na criação das formas de energia que podem existir⁴. É a extensão mínima em que conseguimos

contemplar a “existência” da realidade física. Esta escala minúscula é o novo limiar que todos os atuais candidatos a “nova teoria de tudo” querem alcançar. A teoria das cordas, a teoria M, a geometria não-comutativa, a gravidade quântica em *loop*, os *twisters*... Todos procuram uma nova forma de descrever o que acontece ao nosso papel quando é dividido ao meio 114 vezes.

O que é que acontece se duplicarmos o tamanho da nossa folha de papel A4, passando para o A3, depois para o A2 e por aí em diante? Ao fim de apenas 90 duplicações passámos todas as estrelas e galáxias visíveis e alcançámos o limiar de todo o Universo visível, a 14 mil milhões de anos-luz de distância. Não há dúvida de que há muito mais Universo para lá deste limite, mas esta é a maior distância que a luz teve tempo de percorrer até chegar a nós, desde que a expansão do Universo se iniciou há 14 mil milhões de anos. É o nosso horizonte cósmico.

Ao aproximar o grande e o pequeno, descobrimos que apenas 204 divisões ou duplicações do tamanho de uma folha de papel nos levam à maior e menor dimensão da realidade física, das origens quânticas do espaço aos limites do Universo visível.

Encontrar uma ocorrência difícil deste famoso problema difícil pode ser um problema difícil de resolver.

Brian Hayes

É preciso muito tempo para fazer um *puzzle*, mas apenas um instante para verificar se o problema do *puzzle* foi resolvido. O nosso computador demora uma fração de segundo a multiplicar dois números grandes, mas o leitor (com o seu computador) demoraria muito tempo a encontrar os dois fatores que foram multiplicados para formar um número grande. Há muito que se suspeita, mas nunca foi provado nem refutado (e há um prémio de um milhão de dólares para quem consiga fazer uma das duas coisas), que existe uma divisão real entre problemas “difíceis” e “fáceis” que reflete a quantidade de tempo que é necessário despendar em cálculos para os resolver.

A maioria dos cálculos, ou das tarefas de recolha de informações, que temos de fazer manualmente, como preencher as nossas declarações de impostos, têm em comum o facto de a quantidade de cálculos necessários aumentar em proporção com o número de elementos que temos de processar. Se tivermos três fontes de rendimento, temos o triplo do trabalho. Da mesma forma, o nosso computador demora três vezes mais tempo a fazer o *download* de um ficheiro três vezes maior. Dez livros demorarão, geralmente, dez vezes mais tempo a ler do que apenas um. Este padrão é característico dos problemas “fáceis”. Podem não ser fáceis no sentido habitual do termo, mas se somarmos muitos deles, o tempo necessário para os resolver não aumenta muito depressa. Os computadores conseguem resolver facilmente estes problemas.

Infelizmente, encontramos com frequência outro tipo de problemas que é muito mais difícil de controlar. De cada vez que adicionamos um

elemento extra ao nosso cálculo, verificamos que o tempo necessário para resolver o problema *duplica*. Em breve, o tempo total necessário torna-se extraordinariamente longo e mesmo os computadores mais rápidos da Terra podem ser facilmente derrotados. Estes são aqueles a que chamamos problemas “difíceis”⁵.

Surpreendentemente, os problemas difíceis não têm de ser horrivelmente complicados nem incrivelmente complexos. Simplesmente envolvem muitas possibilidades. Multiplicar dois números primos grandes é uma tarefa computacional “fácil”. Consegue fazer este cálculo de cabeça, com um lápis e papel ou numa calculadora, como preferir. Mas se der o resultado a outra pessoa e lhe pedir que encontre os dois números primos que foram usados na multiplicação, a pessoa em questão poderá defrontar-se com uma vida inteira de investigação usando os computadores mais rápidos do mundo.

Se quiser experimentar resolver um destes problemas “difíceis”, um que parece enganadoramente fácil, tente encontrar os dois números primos que somados dão 389 965 026 819 938*.

Estas operações “alçapão” – assim chamadas porque, tal como quando se desce por um alçapão, é muito mais fácil ir numa direção do que na outra – não são necessariamente uma coisa má. Dificultam-nos a vida, mas também dificultam a vida das pessoas a quem queremos fazer a vida difícil por um muito bom motivo. Todos os principais códigos de segurança do mundo exploram operações alçapão. Sempre que faz compras na Internet e levanta dinheiro numa máquina multibanco, usa estas operações. O seu número PIN é combinado com grandes números primos de forma a obrigar qualquer *hacker* ou criminoso informático que esteja a tentar roubar os seus dados bancários a ter de encontrar dois grandes números primos que foram multiplicados para conseguir obter essa informação. Esta operação não é impossível na sua essência, mas na prática é impossível de realizar num período razoável. Um criminoso com o computador mais rápido do mundo em seu poder poderia conseguir decifrar a encriptação ao fim de vários anos, mas nessa altura os códigos e os números de conta já teriam mudado.

Por este motivo, os grandes números primos são uma coisa muito valiosa e alguns até foram patenteados quando são escritos de determinada

* A resposta é 5569 + 389 965 026 814 369.

forma. Não há um limite de números primos – são infinitos –, mas existe um que é considerado o maior que foi possível verificar, garantindo que não tem fatores. Não existe uma fórmula mágica que consiga gerar todos os números primos e suspeita-se de que pode não vir a existir. Se fosse possível, e se essa fórmula fosse encontrada, haveria uma crise global. Qualquer departamento governamental que a encontrasse iria, sem dúvida, mantê-la em segredo. Um acadêmico que a encontrasse e a tornasse pública sem aviso prévio destruiria o mundo inteiro. Todos os códigos militares, diplomáticos e bancários tornar-se-iam facilmente decifráveis de um dia para o outro usando um computador. A continuidade do comércio virtual ver-se-ia gravemente ameaçada. Teríamos de mudar para sistemas baseados no reconhecimento da íris ou das impressões digitais ou do ADN que utilizam traços únicos da nossa bioquímica em vez de números armazenados na nossa memória. Mas mesmo estes novos indicadores teriam de ser armazenados de uma forma segura.

A fatorização de números primos é um problema “difícil”. Mesmo que seja decifrado e acabe por se revelar um problema “fácil” por meio de uma fórmula mágica, poderíamos usar outro problema “difícil” para encriptar as informações confidenciais, de modo a que demore demasiado tempo a inverter a operação para extrair essas informações. No entanto, sabe-se que, se um dos problemas que consideramos “difíceis” vier a revelar-se fácil por meio de uma nova descoberta, essa descoberta pode ser usada para transformar todos os problemas “difíceis” em problemas “fáceis”. Seria um acontecimento histórico.

Sempre pensei que o recorde se manteria até ser ultrapassado.

Yogi Berra

Os matemáticos sempre se interessaram por recordes: os maiores, os menores e os mais quentes. Serão eles, de alguma forma, previsíveis?

A princípio, poderá pensar que não. É verdade que têm a tendência para “melhorar” – não seriam recordes se assim não fosse –, mas como é que se podia prever que o Michael Johnson ou o Ian Thorpe iam aparecer e bater recorde após recorde? Surpreendentemente, o recorde feminino no salto à vara foi ultrapassado em oito ocasiões diferentes por Yelena Isinbayeva num único ano. Os recordes deste tipo não são aleatórios num sentido muito importante. Cada novo recorde resulta de um esforço competitivo que não é independente de todas as tentativas anteriores de o alcançar. Os praticantes de salto à vara aprendem novos pormenores técnicos e treinam continuamente para vencerem os seus pontos fracos e melhorarem a sua técnica. A única coisa que podemos prever em recordes deste tipo é que serão estabelecidos novamente, com o passar do tempo, embora possa transcorrer muito tempo até surgir um novo.

No entanto, há diferentes tipos de recordes que surgem em sequências de acontecimentos e que se supõe serem independentes entre si. Bons exemplos são os níveis recorde de precipitação mensal, temperaturas recorde, sejam máximas ou mínimas, num lugar ao longo de centenas de anos ou alturas recorde das marés. A suposição de que cada acontecimento é independente do anterior é muito poderosa e permite-nos fazer uma previsão espantosa da probabilidade de haver um novo recorde – independentemente da área a que se aplique. Pode ser de precipitação, queda de neve, queda de folhas, nível das águas, velocidades do vento ou temperatura.

Tomemos como exemplo os níveis de precipitação anuais no Reino Unido. No primeiro ano em que mantemos um registo, a precipitação obtida será inevitavelmente um recorde. No segundo ano, os níveis de precipitação são independentes dos do ano 1, portanto há uma probabilidade de $1/2$ de haver um recorde que ultrapasse os níveis de precipitação do ano 1 e uma probabilidade de $1/2$ de os níveis deste ano não ultrapassarem os do ano 1. Assim, o número de anos recorde que podemos esperar nos primeiros dois anos é $1 + 1/2$. No ano 3 há apenas duas formas de se registar um recorde entre as seis possibilidades dos anos anteriores (ou seja, uma probabilidade de 1 em 3). Desta forma, o número de anos recorde esperado ao fim de 3 anos a manter um registo é de $1 + 1/2 + 1/3$. Se continuar e aplicar o mesmo raciocínio a cada novo ano, descobrirá que ao fim de n anos independentes a recolher dados, o número esperado de anos recorde é a soma de uma série de n frações:

$$1 + 1/2 + 1/3 + 1/4 + \dots + 1/n$$

Esta é a famosa série a que os matemáticos chamam série “harmónica”. Vamos designar a soma desta por $H(n)$ depois de termos somado n termos; desta forma, vemos que $H(1) = 1$, $H(2) = 1,5$, $H(3) = 1,833$, $H(4) = 2,083$, e por aí fora. O mais interessante acerca da soma desta série é que cresce muito lentamente à medida que o número de termos aumenta*, uma vez que $H(256) = 6,12$, mas $H(1000)$ é apenas igual a 7,49 e $H(1\,000\,000) = 14,39$.

Que conclusão podemos tirar disto? Suponhamos que queríamos aplicar a nossa fórmula aos valores recorde de precipitação num dado local no Reino Unido de 1748 a 2004 – um período de 256 anos. Nesse caso, prevemos que deveríamos encontrar apenas $H(256) = 6,12$ ou cerca de 6 anos recorde de precipitação elevada (ou reduzida). Se analisarmos os valores de precipitação registados pelo jardim botânico Kew Gardens para este período, constatamos que foi exatamente o número de anos recorde que houve. Teríamos de esperar mais de 1000 anos para termos uma boa probabilidade de encontrar apenas 8 anos recorde. Os recordes são muito raros quando se referem a acontecimentos aleatórios.

* De facto, quando n se torna muito grande $H(n)$ aumenta apenas à mesma velocidade que o logaritmo de n e alcança um valor aproximado de $0,58 + \ln(n)$.

No passado recente tem vindo a surgir uma preocupação crescente por todo o mundo em relação às evidências de mudanças sistemáticas no clima, o dito “aquecimento global”, e observou-se um número preocupante de recordes climáticos em diferentes lugares. Se os novos recordes se tornarem mais frequentes do que a série harmónica prevê, isso significa que os acontecimentos climáticos já não são acontecimentos anuais independentes, mas que estão a tornar-se parte de uma tendência sistemática e não-aleatória.

29 | “LOTARIA FAÇA-VOCÊ-MESMO”

A Lei Forte dos Números Pequenos: Não existem números pequenos em quantidade suficiente para satisfazer a procura.

Richard Guy

Se precisar de um jogo simples mas intrigante para manter os seus convidados entretidos durante algum tempo, pode experimentar fazer aquilo a que chamei “Lotaria Faça-Você-Mesmo”. Peça a cada uma das pessoas que escolha um número inteiro positivo e que o escreva num cartão juntamente com o seu nome. O objetivo é escolher o número mais pequeno *que não tenha sido escolhido por mais ninguém*. Haverá uma estratégia para ganhar? Poderá pensar que a melhor opção é escolher os números mais pequenos, como 1 ou 2. Mas as outras pessoas poderão ter pensado o mesmo e assim acabará com um número que já foi escolhido por outra pessoa e perderá. Se escolher um número muito grande – há uma quantidade infinita deles à escolha –, perderá certamente. É demasiado fácil alguém escolher um número mais pequeno que o seu. Isto sugere que os melhores números são os que estão entre estes dois extremos. Mas onde? Que tal o 7 ou o 11? De certeza que mais ninguém se vai lembrar de escolher o 7, não é?

Não sei se existe uma estratégia para ganhar, mas o que o jogo evidencia é a nossa relutância em nos considerarmos “típicos”. Sentimo-nos tentados a pensar que poderíamos escolher um número baixo por um motivo de que mais ninguém se lembraria. Claro que a razão pela qual as sondagens de opinião conseguem prever em quem vamos votar, o que vamos comprar, onde vamos passar as nossas férias e a forma como reagimos a um aumento das taxas de juro é, precisamente, o facto de sermos tão semelhantes.

Tenho outra suspeita em relação a este jogo: embora exista uma seleção infinita de números à nossa escolha, esquecemo-nos da maioria deles. Estabelecemos o nosso horizonte aproximadamente no 20, ou no correspondente ao dobro do número de participantes no jogo, se forem mais, e achamos que ninguém vai escolher um número superior a este. Depois excluímos os primeiros cinco números segundo a teoria de que são uma escolha demasiado óbvia e escolhemos entre os que restam com uma probabilidade aproximadamente igual.

Um estudo sistemático das preferências envolveria jogar este jogo muitas vezes com uma grande amostra de jogadores (digamos 100 para cada experiência), de modo a podermos considerar o padrão dos números escolhidos e as escolhas vencedoras. Também seria interessante ver como os jogadores alteram a sua estratégia se participarem várias vezes no jogo. Uma simulação por computador não seria necessariamente útil porque teria de receber instruções sobre a estratégia a adotar. Claramente, a escolha dos números não é aleatória (pois nesse caso a probabilidade de cada número ser escolhido seria igual). A psicologia é importante. Cada jogador tenta imaginar que número os outros jogadores vão escolher. No entanto, a tentação de considerar que não pensamos de forma igual às outras pessoas é tão grande que quase todos caímos nesse erro. Claro que se houvesse uma estratégia infalível para escolher o número mais baixo, todos os jogadores estariam a agir de forma lógica de modo a adotá-la, mas essa estratégia impedi-los-ia de escolher um número que não fosse escolhido por mais ninguém e, conseqüentemente, nunca conseguiriam ganhar.

Será preciso fumar três cachimbos para resolver esse problema, portanto peço-lhe que não fale comigo durante cinquenta minutos.

Sherlock Holmes

Suponhamos que o leitor aparece num concurso televisivo. O apresentador transbordante de entusiasmo mostra-lhe três caixas, com os rótulos A, B e C. Uma delas contém um cheque em seu nome no valor de um milhão de euros. As outras duas contém fotografias do apresentador. Ele sabe qual das caixas contém o cheque que será seu se escolher a caixa certa. O leitor opta pela caixa A. O apresentador começa por abrir a caixa C e mostra a toda a gente que contém uma fotografia dele. O cheque deve estar na caixa A ou na caixa B. O leitor tinha selecionado a caixa A. Agora, o apresentador pergunta-lhe se quer manter a sua escolha original – a caixa A – ou se quer mudar para a caixa B. O que deve fazer? Talvez sinta o impulso de mudar a sua escolha e optar antes pela caixa B, ao mesmo tempo que outra voz na sua cabeça lhe diz: “Fica com a caixa A, ele está apenas a tentar convencer-te a escolher a opção mais barata para o patrão dele.” Ou talvez uma voz mais racional esteja a dizer-lhe que não faz a menor diferença, porque o cheque continua no lugar onde estava inicialmente e a sua escolha inicial tanto pode estar certa como errada.

A resposta é surpreendente; deve mudar a sua escolha para a caixa B! Se o fizer, irá *duplicar* as suas hipóteses de acertar na caixa que contém o cheque. Se mantiver a escolha inicial da caixa A, terá apenas 1/3 de probabilidades de ganhar o cheque; mas se mudar para a caixa B, as suas probabilidades aumentam para 2/3.

Como é que isto é possível? Inicialmente, há uma probabilidade de 1 em 3 de o cheque estar em qualquer uma das caixas. Isto significa que há 1/3 de probabilidades de estar na caixa A e 2/3 de probabilidades de estar na

B ou na C. Quando o apresentador intervém e abre uma caixa, isto não altera as probabilidades, porque ele vai sempre escolher uma das caixas que não contém o cheque. Assim, depois de ele abrir a caixa C, ainda há $1/3$ de probabilidades de o cheque estar na caixa A, mas agora há $2/3$ de probabilidades de estar na caixa B, porque já não temos dúvidas de que não está na C, portanto, deve mudar.

Ainda não está convencido? Vejamos as coisas por outro ponto de vista. Depois de o apresentador abrir a caixa C, o leitor tem duas opções: pode manter a sua escolha inicial e optar pela caixa A, o que garantirá que ganha se o seu palpite inicial estiver certo. Em alternativa, pode mudar para a caixa B e nesse caso só ganhará se o seu palpite inicial estiver errado. A sua primeira escolha, a caixa A, estará certa $1/3$ das vezes e errada $2/3$ das vezes. Desta forma, ao mudar de opinião ganhará o cheque $2/3$ das vezes ao passo que manter o palpite original só será uma estratégia vencedora $1/3$ das vezes.

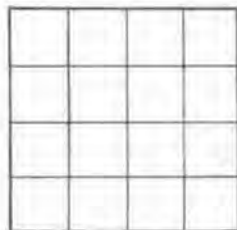
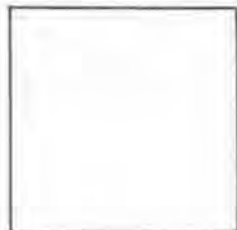
Espero tê-lo convencido com esta experiência revolucionária.

Vou mostrar-te o medo num punhado de poeira.

T. S. Elliot, *A Terra Devastada*

Uma das lições que aprendemos com vários incêndios catastróficos é que o pó é mortal. Um pequeno incêndio num velho armazém pode alastrar e transformar-se num inferno explosivo devido aos esforços para o extinguir se esses esforços espalharem grandes quantidades de pó no ar, onde entra em combustão e espalha o fogo num instante. Qualquer espaço escuro e abandonado – vãos de escadas, ou de bancadas de um estádio ou espaços de armazenamento negligenciados – onde o pó pode acumular-se em grandes quantidades e não ser detetado, constitui um grave risco de incêndio.

Porquê? Normalmente não pensamos no pó como sendo uma matéria especialmente inflamável. O que é que o transforma numa presença tão perigosa? A resposta está relacionada com a geometria. Comece com um quadrado de tecido e corte-o em 16 quadrados menores. Se as dimensões do quadrado original fossem 4 cm x 4 cm, cada um dos quadrados menores terá 1 cm x 1 cm. A área total da superfície do tecido é a mesma – 16 centímetros quadrados. Nada se perdeu. A grande mudança, contudo, é a dimensão das arestas expostas. O quadrado original tinha um perímetro com o comprimento de 16 cm, mas os quadrados mais pequenos têm um perímetro de 4 cm cada e são 16, portanto o perímetro total ficou com o quádruplo do tamanho: passou para $4 \times 16 \text{ cm} = 64 \text{ cm}$.



Se fizermos o mesmo com um cubo, este terá 6 faces (como um dado) de 4 cm x 4 cm e cada uma delas terá uma área de $16 \text{ cm}^2 = 96 \text{ cm}^2$. No entanto, se tivéssemos dividido este cubo grande em 64 cubos mais pequenos, cada um com as dimensões 1 cm x 1 cm x 1 cm, o volume total manter-se-ia, mas a área de superfície total de todos os cubos mais pequenos (cada um deles com 6 faces com uma área de 1 cm x 1 cm) teria aumentado para $64 \times 6 \times 1 \text{ cm}^2$.

O que estes exemplos simples demonstram é que, se alguma coisa for dividida em partes menores, a superfície total dos fragmentos aumenta muito ao mesmo tempo que estes diminuem. O fogo alimenta-se das superfícies, pois é nelas que os materiais combustíveis podem entrar em contacto com o ar no qual está contido o oxigénio de que o fogo necessita para se manter ativo. É por isso que rasgamos pedaços de papel quando queremos atear uma fogueira. Um bloco único de material combustível arde lentamente por apenas uma pequena parte dele estar em contacto com o ar circundante e é nesse contacto com o ar que se dá a combustão. Se esse pedaço se desfizer numa poeira de pequenos fragmentos, haverá muito mais área de superfície do material em contacto com o ar e a combustão dá-se em toda a parte, espalhando-se rapidamente de um fragmento de pó para o seguinte. O resultado pode ser um incêndio súbito ou uma explosão quando a densidade de partículas de pó no ar for tão grande que todo o ar serve de combustível às chamas.

De um modo geral, muitas coisas pequenas representam sempre um risco de incêndio maior do que uma coisa grande com o mesmo volume e composição. A exploração florestal desorganizada, em que se cortam todas as árvores grandes, deixando hectares de lascas de madeira e serradura espalhadas pelo chão da floresta, é um exemplo atual pertinente.

As poeiras são perigosas em grandes quantidades. Na década de 1980 aconteceu na Grã-Bretanha uma catástrofe, quando uma grande fábrica nas Midlands onde era produzido pudim em pó se incendiou. Basta deixar uma pitada de leite em pó, farinha ou serradura numa pequena chama para produzir uma autêntica labareda, com vários metros de altura. (*Não experimente fazer isto em casa! Veja antes o vídeo*.*)

* Para ver uma sequência de imagens de uma demonstração feita por Neil Dixon para a aula de química na escola, visite o website: <http://observer.guardian.co.uk/flash/page/0,1927850,00.html>.

A principal causa dos problemas são as soluções.

Regra de Severeid

Há um problema clássico que consiste em fazer uma escolha entre um grande número de candidatos; talvez seja um gerente que se depare com a escolha entre 500 candidatas ao cargo de secretária da empresa, ou um rei que tenha de escolher uma esposa de entre todas as jovens do seu reino, ou uma universidade que precise de escolher o melhor aluno a admitir de uma longa lista de candidatos. Quando o número de candidatos é moderado, podemos entrevistá-los a todos, avaliar cada um em comparação com todos os outros, voltar a entrevistar aqueles em relação aos quais ainda estamos indecisos e escolher o melhor para a função. Se o número de candidatos for muito grande, essa pode não ser a linha de ação mais prática. Pode escolher aleatoriamente um candidato, mas se houver N candidatos, a escolha de encontrar o melhor através de uma seleção aleatória é apenas de $1/N$, e sendo N um número grande, essa probabilidade torna-se muito reduzida – menos de 1% nos casos em que haja mais de 100 candidatos. Para encontrar o melhor candidato, o primeiro método referido, o de entrevistar todos os candidatos, consome demasiado tempo, mas é fiável; a alternativa da escolha aleatória é rápida, mas muito pouco fiável. Haverá um método melhor, algures entre estes dois extremos, que nos dê boas probabilidades de encontrar o melhor candidato sem despende uma quantidade excessiva de tempo?

Existe, e a sua simplicidade e eficácia relativa são duplamente surpreendentes, portanto vamos estabelecer as regras base. Temos N candidatos a uma “função” e vamos considerar as suas candidaturas por uma ordem aleatória. Depois de termos analisado uma candidatura, podemos marcar a sua classificação relativa às outras candidaturas já analisadas, embora

só tenhamos interesse em saber qual é o melhor candidato que analisámos até agora. Depois de termos analisado uma candidatura, não poderemos voltar a analisá-la. O nosso único objetivo é encontrar o melhor candidato. Todas as outras escolhas são fracassos. Assim, depois de termos entrevistado qualquer candidato, só precisamos de anotar qual é o melhor de todos os candidatos (incluindo este último) que vimos até agora. Quantas mais candidaturas entre as N iniciais precisamos de analisar para termos a melhor probabilidade de escolher o candidato mais forte e que estratégia devemos adotar?

A nossa estratégia consistirá em entrevistar os primeiros C dos N candidatos da lista e depois escolher o seguinte dos restantes candidatos que concluirmos que é melhor que todos eles. Mas como é que vamos determinar o número C ? É essa a questão.

Imagine que temos 3 candidatos: 1, 2 e 3, e que o número 3 é melhor do que o número 2 que, por sua vez, é melhor do que o número 1; as seis ordens possíveis em que podemos entrevistá-los são:

123 132 213 231 312 321

Se decidíssemos escolher sempre o primeiro candidato que entrevistámos, esta estratégia só escolheria o melhor (o número 3) em dois dos seis padrões de entrevista, pelo que escolheríamos o melhor candidato com uma probabilidade de $2/6$ ou $1/3$. Se rejeitássemos sempre o primeiro candidato e escolhêssemos o seguinte candidato entrevistado com a melhor classificação, escolheríamos o melhor candidato apenas no segundo (132), terceiro (213) e quarto (231) casos, pelo que a possibilidade de encontrar o melhor candidato é agora de $3/6$ ou $1/2$. Se rejeitássemos dois candidatos e escolhêssemos o terceiro que entrevistássemos com a melhor classificação, obteríamos o melhor candidato apenas no primeiro (123) e no terceiro casos (213), sendo a nossa probabilidade de escolher o melhor novamente de apenas $1/3$. Desta forma, num cenário em que temos 3 candidatos, a estratégia de rejeitar o primeiro e escolher o seguinte com a melhor classificação dá-nos as melhores probabilidades de selecionar o melhor candidato.

Este tipo de análise pode estender-se a uma situação em que o número de candidatos, N , é superior a 3. Com 4 candidatos, há 24 ordens possíveis pelas quais podemos entrevistá-los a todos. Neste caso, descobrimos que

a estratégia de deixar passar o primeiro candidato e escolher o seguinte com a melhor classificação continua a dar-nos a melhor probabilidade de encontrar o melhor, e fá-lo com uma taxa de sucesso* igual a $11/24$. Este argumento pode ser aplicado a qualquer número de candidatos.

À medida que o número de candidatos aumenta, a estratégia e o resultado aproximam-se cada vez mais daquele que pode ser considerado ideal. Consideremos um caso em que existem 100 candidatos. A estratégia ideal⁶ consiste em entrevistar 37 deles e escolher o seguinte que encontrarmos que seja melhor do que todos os anteriores e não entrevistar mais ninguém. Esta estratégia terá como resultado a escolha do melhor candidato para a função com uma probabilidade de cerca de 37,1% – bastante boa em comparação com a probabilidade de 1% que teríamos se tivéssemos escolhido aleatoriamente⁷.

Este tipo de estratégia deveria ser usado na prática? Faz muito sentido dizer que no contexto de uma entrevista de emprego se deve entrevistar todos os candidatos, mas e se aplicarmos o mesmo raciocínio a um processo de recrutamento de novos executivos, à sua “busca” da mulher da sua vida, do projeto do próximo *bestseller* ou do lugar perfeito para viver? Não pode procurar durante toda a sua vida. Quando é que deve parar de procurar e tomar uma decisão? Ou um cenário menos crítico, se estiver a procurar um hotel para se hospedar ou um restaurante onde comer, ou a pesquisar os melhores preços de uma viagem na Internet ou a bomba de gasolina com os melhores preços? Quantas opções deve considerar antes de tomar uma decisão? Todos estes são problemas sequenciais do tipo que estivemos a analisar na nossa busca de uma estratégia ideal. A experiência sugere que não procuramos o suficiente antes de fazermos a nossa escolha definitiva. Devido às pressões psicológicas ou simplesmente por impaciência (nossa ou dos outros), somos forçados a fazer uma escolha antes de termos analisado uma fração crítica, 37%, das opções.

* Escolher o primeiro ou o último candidato dá-nos sempre uma probabilidade de $1/4$, deixar passar dois candidatos dá-nos uma probabilidade de $5/12$. Deixar passar apenas um candidato dá-nos uma probabilidade de $11/24$, que é a solução ideal.

33 | DIVÓRCIO AMIGÁVEL: A SOLUÇÃO VANTAJOSA PARA AMBAS AS PARTES

Conrad Wilson foi muito generoso comigo no divórcio. Deu-me 5000 Bíblias de Gideon*.

Zsa Zsa Gabor

“Partilhar é uma coisa boa, Pai”, disse o nosso filho de três anos uma vez ao olhar para o meu gelado depois de ter acabado de comer o dele. Mas partilhar não é assim tão simples. Se precisar de dividir algo por duas ou mais pessoas, o que deve tentar fazer? É fácil pensar que basta fazer a divisão que *considera* justa e, no caso das duas pessoas, isto significa dividir o dito bem ao meio. Infelizmente, embora esta estratégia possa funcionar quando se divide algo que é muito simples, como uma soma de dinheiro, não é uma estratégia óbvia quando o bem que está a ser dividido significa coisas diferentes para pessoas diferentes. Se precisarmos de dividir uma extensão de terra entre dois países, cada um deles pode valorizar de forma diferente um determinado aspeto, como água para a agricultura ou montanhas para o turismo. Alternativamente, as coisas que estão a ser divididas podem envolver aspetos que nenhuma das partes deseja, como tarefas domésticas ou ir para a fila de uma repartição.

No caso de um divórcio, há muitas coisas que podem ser partilhadas, mas cada pessoa atribui um valor diferente a coisas diferentes. Uma pessoa pode dar mais valor à casa e a outra à coleção de quadros ou ao cão da família. Embora o leitor, como possível mediador, tenha uma perspectiva

* A organização cristã Gideon's International (também conhecida como Gideon's Bible (sendo os seus membros apelidados de apenas "Gideon"), dedica-se a distribuir gratuitamente exemplares da Bíblia Sagrada, sobretudo deixando-os em quartos de hotel. (N. R.)

própria do valor dos diferentes bens a ser partilhados, as duas partes atribuem valores diferentes às diversas partes constituintes da propriedade comum. O objetivo de um mediador deverá ser chegar a uma divisão que agrade a ambas as partes. Isto não significa que as metades têm de ser "iguais" em nenhum sentido numérico.

Uma forma simples e tradicional de agir é pedir a um dos membros do casal que especifique uma divisão dos bens em duas partes e depois pedir ao outro que escolha qual das partes quer para si. A pessoa que especifica a divisão tem um incentivo para ser escrupulosamente justa, porque pode acabar por receber a parte menos justa se a outra escolher a "melhor" metade. Este método deverá conseguir evitar toda a inveja que possa existir no processo de divisão (a menos que a pessoa que faz a divisão saiba algo acerca dos bens que a outra não sabe – por exemplo, que existem reservas de petróleo debaixo de uma determinada extensão de terreno). No entanto, há um potencial problema. As duas pessoas podem, ainda assim, valorizar de forma diferente as diversas partes, pelo que o que parece bom aos olhos de um poderá não ser bom para o outro.

Steven Brams, Michael Jones e Christian Klamler sugeriram uma forma melhor de dividir os bens entre duas partes que ambas considerarão justa. Pede-se a cada parte que diga ao mediador como procederia para dividir os bens de forma justa. Se ambos fizerem uma escolha idêntica, não há problema e acordam imediatamente o que se deve fazer. Se não concordarem, o mediador terá de intervir.

A B

Suponhamos que os bens são distribuídos ao longo de uma linha e que a minha escolha de uma divisão justa divide a linha no ponto A, mas a sua divide-a no ponto B. Desta forma, a divisão justa atribui-me a parte à esquerda do ponto A e dá ao leitor a parte à direita do ponto B. Entre os dois existe uma porção remanescente que o mediador divide ao meio, dando metade a cada um. Neste processo, ambos acabamos a receber mais do que a "metade" que tínhamos esperado. Ambos ficamos satisfeitos.

Poderíamos fazer ainda melhor do que a Brahms and Co. sugere se não nos limitássemos a pedir ao mediador que dividisse o remanescente por dois. Poderíamos repetir o processo de divisão justa para essa porção,

devendo cada um escolher qual considera que seria a divisão justa, excluindo as nossas duas porções que não se sobrepõem de modo a que restasse uma nova porção (agora menor) e depois dividi-la da mesma forma, e por aí em diante, até nos restar uma última porção que (segundo acordo prévio) seria considerada negligenciavelmente pequena, ou até as escolhas da divisão do remanescente por cada uma das partes serem iguais.

Se houver três ou mais partes que desejem dividir de forma justa os bens, o processo torna-se muito mais complicado, mas é, na sua essência, o mesmo. As soluções para estes problemas foram patenteadas pela Universidade de Nova Iorque, para poderem ser usadas comercialmente em casos em que seja necessário resolver as disputas e encontrar uma divisão justa dos bens. As suas aplicações vão desde divórcios até ao processo de paz no Médio Oriente.

Ainda és tão jovem que pensas que um mês é muito tempo.

Henning Mankel

Se convidar muitas pessoas para a sua festa de aniversário, poderá ter interesse em saber quantos convidados são necessários para ter mais de 50% de probabilidades de um deles fazer anos no mesmo dia. Suponhamos que não sabe nada acerca das datas de aniversário dos seus convidados na altura do convite. Depois, esquecendo os anos bissextos e assumindo que um ano tem 365 dias, precisará de pelo menos⁸ 253 convidados para ter uma probabilidade superior a 50% de partilhar o aniversário com um deles. O número é muito maior do que 365 a dividir por dois porque é provável que muitos dos convidados tenham a mesma data de aniversário que outros presentes para além de si. Ficaria muito caro encomendar comida para uma festa assim.

Um jogo ainda mais interessante consiste simplesmente em procurar pessoas que partilham a data de aniversário entre elas e não necessariamente consigo. Quantos convidados seriam necessários para ter uma probabilidade superior a 50% de dois deles terem a mesma data de aniversário? Se experimentar fazer esta pergunta a pessoas que não analisaram profundamente o problema, elas sobrestimarão o resultado por uma grande margem. A resposta é surpreendente. Com apenas 23 pessoas⁹, há 50,7% de probabilidades de dois convidados terem a mesma data de aniversário; com 22 convidados essa probabilidade passa para 47,6% e com 24 passa para 53,8%⁹.

Consideremos duas equipas de futebol, juntemos-lhes um árbitro e com este conjunto temos uma probabilidade de 1 para 2 de dois deles

⁸ Se tivermos em conta os anos bissextos, haverá uma pequena diferença, mas este número não se altera.

fazerem anos no mesmo dia. Há uma ligação simples com o primeiro problema que considerámos, o de conseguir ter um convidado a partilhar *consigo* a data de aniversário, que exigia 253 convidados para a probabilidade ser superior a 50%. O motivo pelo qual qualquer par exige apenas a presença de 23 pessoas é o facto de haver tantos pares de 23 pessoas – de facto existem* $(23 \times 22)/2 = 253$ pares.

O matemático americano Paul Halmos encontrou uma útil aproximação para este problema com uma forma ligeiramente diferente. Ele demonstrou que se reunirmos um grande grupo de pessoas, número esse a que chamaremos novamente N , precisaremos de fazer uma seleção aleatória de pelo menos $1,18 \times N^{1/2}$ delas para termos uma probabilidade superior a 50% de duas destas pessoas terem a mesma data de aniversário. Se N for igual a 365, a fórmula obtém o resultado 22,544, pelo que precisamos de 23 pessoas.

Uma das suposições que está subjacente a esta análise é a de que existe uma probabilidade igual de os aniversários calharem em qualquer dia do ano. Na prática, isto não é verdade. A conceção é mais provável nas épocas festivas e é pouco provável que os partos planeados por cesariana calhem no dia de Natal ou na noite da passagem de ano. Os desportistas famosos são outro caso interessante – podemos investigar as datas de aniversário dos jogadores de futebol da primeira divisão ou dos membros da equipa de atletismo do Reino Unido. Aqui, suspeito que encontrará um maior número de nascimentos no outono. A razão de ser não tem nada que ver com os signos astrais. No Reino Unido, o ano escolar começa no início de setembro, pelo que as crianças que fazem anos em setembro, outubro e novembro são significativamente mais velhas do que os alunos da mesma turma que fazem anos em junho, julho e agosto numa altura da vida em que um período de 6 a 9 meses representa uma grande diferença de força física e velocidade. Desta forma, as crianças cuja data de aniversário calha no outono têm mais probabilidade de entrar para as equipas desportivas e têm o impulso, a oportunidade e o treino suplementar necessários para terem sucesso na idade adulta. O mesmo é verdade em relação a outras atividades que exigem outros tipos de maturidade.

* Existem 23 para a primeira escolha e depois 22 para cada um deles, ou seja, 23×22 , mas depois dividimos esse número por 2 porque a ordem pela qual o par é encontrado não tem importância (isto é, "tu e eu" é o mesmo que "eu e tu").

Há uma grande variedade de atividades que exigem que facultemos a nossa data de aniversário como parte dos procedimentos de segurança. Os *websites* dos bancos, das lojas *online* e das companhias aéreas usam as datas de aniversário como parte das verificações de segurança. Já vimos que considerar apenas a data de aniversário não é boa ideia. A probabilidade de dois clientes terem a mesma data de aniversário é muito elevada. Podemos reduzir essas probabilidades pedindo a indicação do ano de nascimento e acrescentando uma palavra-passe. Poderá constatar que o problema é essencialmente o mesmo, com a diferença que, em vez de ter 365 datas que podem ser partilhadas por duas pessoas na sua festa, procura agora a probabilidade de partilharem essa data e, por exemplo, a escolha de uma sequência de dez letras como palavra-passe. Isto reduz dramaticamente as probabilidades de correspondências aleatórias. Se uma palavra-passe for composta por oito caracteres alfabéticos, haverá 26^{10} escolhas possíveis e a fórmula de Halmos levar-nos-ia a esperar necessitar de $1,28 \times 26^5$, ou aproximadamente 15 208 161 clientes para haver uma probabilidade de 50% de dois deles partilharem estas duas informações. Em julho de 2007 estimava-se que a população mundial era de 6 602 224 175 habitantes, pelo que uma palavra-passe composta por uma sequência de 14 letras levava a que o número de pessoas necessário para uma probabilidade de 50% fosse superior à população mundial.

35 | A INCLINAÇÃO DOS MOINHOS DE VENTO

A resposta, meu amigo, levá-a o vento.

Bob Dylan, "Blowing in the wind"

Se viajar pelo Reino Unido encontrará cada vez mais moinhos de vento modernos espalhados pelas zonas rurais, como naves de extraterrestres. A sua presença é controversa. Estão ali colocados para combater a poluição atmosférica, oferecendo uma fonte energética mais ecológica; no entanto, produzem poluição visual quando são colocados de forma inadequada em paisagens rurais ou marítimas imaculadas (mas sem dúvida ventosas).

Há algumas questões interessantes acerca dos moinhos de vento, ou "turbinas eólicas", como são atualmente chamados. Os antigos moinhos de vento tinham quatro velas cruzadas ao centro em forma de X. Os moinhos de vento modernos parecem hélices de aviões e têm habitualmente três pás. Há vários fatores que levaram a que o moinho de vento de três pás (ou de estilo dinamarquês) se tivesse tornado tão comum. Três velas são mais baratas do que quatro. Então, por que não usar apenas duas? Os moinhos de quatro velas tinham uma estranha propriedade que os tornava menos estáveis do que os de três. Os moinhos de vento com quatro velas (ou qualquer outro número par) têm a seguinte característica: quando uma delas está na posição vertical mais elevada, extraíndo a potência máxima do vento, a outra extremidade dessa vela estará apontada verticalmente para baixo, protegida do vento pelo suporte do moinho. Isto causa tensão nas velas e uma tendência para o moinho de vento balançar, o que pode ser perigoso quando se fazem sentir ventos fortes. Os moinhos de vento com três pás (ou qualquer número ímpar de pás) não têm este problema. As três pás estão separadas por uma distância de 120 graus e quando uma delas

está na vertical, nenhuma das outras duas pode estar também na vertical. É certo que três pás apanham menos vento do que quatro, pelo que têm de girar mais depressa para gerar a mesma quantidade de energia.

A eficiência dos moinhos de vento é uma questão interessante, que foi resolvida pela primeira vez em 1919 pelo engenheiro alemão Albert Betz. As velas de um moinho de vento são como um rotor que varre uma certa superfície, A , do ar que chega até ele à velocidade U e sai a outra velocidade, inferior, V . A perda de velocidade do ar como consequência da ação das velas é o que permite ao moinho de vento extrair energia do ar em movimento. A velocidade média do ar a passar nas velas é de $1/2 (U + V)$. A massa de ar por unidade de tempo que passa pelos braços em rotação é $\dot{M} = \rho A \times 1/2 (U + V)$, em que ρ é a densidade do ar. A energia gerada pelo moinho de vento é apenas a diferença entre a velocidade de mudança da energia cinética do ar antes e depois de passar pelos rotores. Esta corresponde a $P = 1/2 \dot{M} U^2$. Se usarmos a fórmula de $\dot{M}$ nesta equação, verificamos que a energia gerada é:

$$P = 1/4 \rho A (U^2 - V^2)(U + V)$$

Porém, se o moinho de vento não tivesse estado ali, a força total do vento não perturbado teria sido apenas de $P_0 = 1/2 \rho A U^3$. Desta forma, a eficiência do moinho de vento na extração de energia da massa móvel de ar será P/P_0 , que é igual a 1, indicando 100% de eficiência como gerador de energia quando $P = P_0$. Se dividirmos as duas últimas fórmulas teremos:

$$P/P_0 = 1/2 \{1 - (V/U)^2\} \times \{1 + (V/U)\}$$

É uma fórmula interessante. Quando V/U é pequeno, P/P_0 é próximo de 1/2; quando V/U se aproxima do seu valor máximo de 1, P/P_0 tende para zero porque não foi extraída nenhuma velocidade do vento. Algures num ponto intermédio, quando $V/U = 1/3$, P/P_0 alcança o seu valor máximo. Nesse caso a eficiência energética alcança o seu valor máximo e é igual a 16/27, ou cerca de 59,26%. Esta é a Lei de Betz para a máxima eficiência energética possível de um moinho de vento ou rotor a extrair energia do ar em movimento. A razão pela qual a eficiência é inferior

a 100% é fácil de compreender. Se fosse de 100%, toda a energia recebida do vento teria de ser removida das velas do moinho – o que poderia ser alcançado, por exemplo, dando ao moinho a configuração de um disco sólido de modo a intercetar todo o vento – mas nesse caso os rotores não girariam. A velocidade do vento a jusante (a nossa quantidade V) seria igual a zero porque não poderia passar mais vento pelas velas do moinho.

Na prática, as boas turbinas eólicas conseguem alcançar níveis de eficiência de cerca de 40%. Quando a energia é convertida em eletricidade utilizável, há menos perdas nos rolamentos e linhas de transmissão, e apenas cerca de 20% da energia eólica disponível acabará por ser transformada em eletricidade utilizável.

A potência máxima que pode ser extraída do vento por qualquer vela, rotor ou turbina ocorre quando $V/U = 1/3$, de modo a que $P_{\text{max}} = (8/7) \times D \times A \times U^3$. Se o diâmetro do círculo do rotor for d , a sua área será $A = \pi d^2/4$ e se o moinho de vento trabalhar com um nível de eficiência de 50%, a energia gerada será cerca de $1,0 \times (d/2\text{m})^2 \times (U/1\text{ms}^{-1})^3$ watts.

Posh e Becks não saíram do quarto de hotel por terem ficado confusos com o letreiro "Não Incomodar" pendurado do lado de dentro da porta.

Angus Deayton

A prestidigitação, vista de perto, pode parecer muito misteriosa e torna-se espantosa se o prestidigitador lhe mostrar como é feita. Com que facilidade ele o enganou; como foi incapaz de ver o que estava mesmo debaixo do seu nariz; como era simples. Percebe rapidamente como seria incapaz de averiguar a autenticidade de fenómenos de torção de colheres ou de levitação. Os cientistas são os mais fáceis de enganar: não estão habituados a ter a Natureza a conspirar contra eles. Acreditam que quase tudo o que veem é verdade. Os mágicos não acreditam em nada.

Com esta ideia em mente, quero contar-lhe uma pequena história matemática, adaptada de uma versão de Frank Morgan, que é um exemplo de prestidigitação verbal. O leitor segue cuidadosamente a história, mas algo parece desaparecer a meio do relato – uma simples soma de dinheiro – e tem de descobrir para onde foi; ou se, de facto, alguma vez lá esteve.

Três viajantes chegam a um hotel já tarde numa certa noite, cada um com 10 libras na carteira. Decidem partilhar um quarto grande e o hotel cobra-lhes 30 libras pela estadia de uma noite, pelo que cada um paga 10 libras. Depois de subirem para o quarto, com o pacote a transportar as malas, a rececionista do hotel recebe um e-mail da sede daquela cadeia de hotéis que a informa de que está a decorrer uma promoção e que o preço do quarto passa a ser de 25 libras para os hóspedes que ali fiquem naquela noite. A rececionista, sendo uma pessoa escrupulosamente honesta, manda de imediato o pacote ao quarto dos hóspedes recém-chegados, para lhes entregar uma nota de 5 libras como reembolso. O pacote é menos escrupuloso. Não tinha recebido gorjeta por ter

transportado as malas dos hóspedes e não percebe como vão dividir as 5 libras pelos três, portanto decide guardar 2 libras, a título de “gorjeta”, e dá aos três hóspedes um reembolso de 1 libra a cada. Desta forma, cada um dos hóspedes pagou 9 libras pelo quarto e o pacote tem 2 libras no bolso, o que perfaz um total de 29 libras. Mas eles pagaram 30 – o que é que aconteceu à outra libra?

* Se analisarmos mais atentamente o argumento, constatamos que não falta 1 libra. No final da história, os três hóspedes têm um total de 3 libras, o pacote tem 2 e o hotel tem 25.

37 | INVESTIMENTOS FINANCEIROS COM VIAJANTES NO TEMPO

Não sou adivinho! Sou apenas um vigarista! Não consigo fazer previsões genuínas. Se eu conseguisse prever como esta história ia acabar, teria ficado em casa.

René Goscinny e Albert Uderzo. *Astérix e o Adivinho*

Imagine que existem três civilizações avançadas no Universo que aperfeiçoaram a arte (e a ciência) de viajar no tempo. É importante perceber que a ideia de viajar para o futuro não é, de forma alguma, controversa. O seu acontecimento está previsto nas teorias de Einstein acerca do tempo e do movimento que descrevem com tanta exatidão o mundo que nos rodeia e que é observado de forma rotineira em experiências do campo da física. Se dois gémeos idênticos fossem separados, ficando um na Terra e o outro sendo enviado para uma viagem de ida e volta ao espaço, o gémeo que ia viajar no tempo seria mais jovem do que o irmão que tinha permanecido na Terra quando regressasse para se encontrar com ele neste planeta. O gémeo regressado teria viajado ao futuro do gémeo que ficou na Terra.

Então, as viagens ao futuro são apenas uma questão prática: a de conseguirmos construir um meio de transporte que suporte a tensão e consiga atingir as necessárias velocidades próximas da velocidade da luz para as tornar realidade. As viagens no tempo até ao passado são uma questão completamente diferente. Este é o campo dos paradoxos da alteração do passado, embora muitos deles sejam baseados numa ideia errada*.

Eis a prova observacional de que os viajantes no tempo não estão envolvidos numa atividade económica sistemática no nosso mundo.

* Ver J. D. Barrow. *The Infinite Book* (Vintage), para mais informações.

O facto económico principal que observamos é que as taxas de juro não são iguais a zero. Se fossem positivas, aqueles que viajassem até ao passado poderiam usar o seu conhecimento dos preços das ações conhecidos no futuro para investir naquelas ações cujo valor sabiam que ia aumentar mais. Obteriam grandes lucros nos mercados de investimentos. Consequentemente, as taxas de juro tenderiam para zero. Em alternativa, se as taxas de juro fossem negativas (significando que os investimentos valeriam menos no futuro), os viajantes no tempo poderiam vender os seus investimentos ao elevado preço atual e voltar a comprá-los no futuro a um preço inferior, apenas para voltarem atrás no tempo e tornarem a vendê-los a um preço elevado. Mais uma vez, a única forma de os mercados pararem esta máquina financeira perpétua seria baixando as taxas de juro para zero. Assim, a constatação de que existem taxas de juro que não são de zero significa que este tipo de atividade de especulação financeira não está a ser levada a cabo por viajantes no tempo vindos do futuro*.

O mesmo tipo de argumento seria aplicável aos casinos e a outras formas de jogo organizado. Aliás, estes poderão ser um melhor alvo para os viajantes no tempo porque não pagam impostos! Se souber quem vai ser o futuro vencedor do *derby* ou o próximo número em que a bola da roleta vai parar, uma curta viagem ao passado garantirá o prémio. O facto de ainda existirem casinos e outras formas de jogo – e de terem um lucro significativo – é, mais uma vez, uma forte indicação de que não existem jogadores que viajam no tempo.

Estes exemplos podem parecer bastante rebuscados, mas não é verdade que os mesmos argumentos podem ser aplicados a várias formas de percepção extrassensorial ou de conhecimento parapsíquico? Qualquer pessoa que conseguisse prever o futuro teria uma vantagem tal que conseguiria acumular uma riqueza fabulosa rápida e facilmente. Essa pessoa poderia ganhar a lotaria todas as semanas. Se existissem intuições fiáveis em relação ao futuro em algumas mentes humanas (ou pré-humanas), essas dariam aos respetivos donos uma tremenda vantagem evolucionária. Poderia

* Este argumento foi inicialmente estabelecido pelo economista californiano Marc Reinganum. Para os investidores mais conservadores, é de notar que o valor de 1 libra investida em 2007 numa conta poupança que produzisse um juro composto de 4% teria aumentado para $1\text{£} \times (1 + 0,04)^{100} = 108$ milhões de milhões de milhões no ano de 3007. Contudo, este seria provavelmente o preço do jornal de domingo nessa altura.

prever os perigos e planejar o futuro sem incerteza. Um gene que desse ao seu dono esta apólice de seguro contra todas as eventualidades espalhar-se-ia rapidamente e aqueles que o tivessem acabariam por dominar a população. O facto de as capacidades psíquicas serem aparentemente tão raras é um argumento muito forte contra a sua existência.

38 | UM PENSAMENTO PELOS SEUS TOSTÕES

Um dos malefícios do dinheiro é tentar-nos a olhar para ele em vez de olharmos para as coisas que compra.

E. M. Forster

As moedas causam atrasos. Basta comprar algo que custe 79 *pence* (ou 79 cêntimos) e desatará a procurar no fundo da mala a combinação exata de moedas que perfaça essa quantia. Se não o fizer e optar antes por entregar ao vendedor 1£ (ou 1€), receberá ainda mais trocos e da próxima vez demorará ainda mais tempo a encontrar as moedas. A questão que se coloca é: qual é a combinação de moedas mais adequada para fazer este pagamento?

No Reino Unido (e no resto da Europa) temos moedas com os valores 1, 2, 5, 10, 20 e 50 para somar 1 libra (ou 1 euro). Nos Estados Unidos há moedas de 1, 5, 10 e 25 cêntimos, com as quais se pode somar 1 dólar ou 100 cêntimos. Estes sistemas são “práticos” no sentido em que o menor número de moedas necessário para formar qualquer valor igual ou inferior a 100 é usado seguindo uma receita simples a que os informáticos chamam “algoritmo ganancioso”. O nome deve-se ao facto de esta receita se basear em começar por somar a maior quantidade possível das maiores moedas, depois fazer o mesmo com a segunda maior moeda, e por aí em diante. A melhor maneira de somar 76 cêntimos de dólar utilizaria apenas três moedas: uma moeda de 50 cêntimos, uma de 25 cêntimos e uma de 1 cêntimo. No Reino Unido, para somar 76 *pence* necessitaríamos de uma moeda de 50 *pence*, uma de 20 *pence*, uma de 5 *pence* e uma de 1 *pence*, necessitando de um máximo de quatro moedas – mais uma do que é necessária no sistema americano, apesar de termos mais moedas à escolha.

O meu primeiro ano de escola foi memorável, pelo facto de ter incluído o "Dia D", 15 de fevereiro de 1971, quando o Reino Unido mudou para um sistema monetário decimal. Até àquela data, o sistema antigo, composto por libras, *shillings* e *pence*, tinha um grande número de moedas com valores invulgares (por motivos históricos). Uma moeda de 1 libra valia 240 moedas de 1 *penny* antigas (cujas abreviatura era d, do latim *denarii*) e na minha infância havia moedas com os valores ½d, 1d, 3d, 6d, 12d, 24d e 30d, que tinham como nomes comuns *halfpenny*, *penny*, *threepenny bit*, *sixpence*, *shilling*, *florin* e *half crown*. Este sistema não tinha o carácter prático que torna os cálculos mais simples usando o algoritmo ganancioso. Para pagar 48d (4 *shillings*), o algoritmo dir-nos-ia para usarmos uma moeda de 30d, uma de 12d e uma de 6d – um total de três moedas. Mas o mesmo podia ser feito de forma mais eficaz com duas moedas de 24 *pence*. Nos séculos anteriores houve uma moeda de 4d, chamada *groat*, que tinha o mesmo efeito nos pagamentos de baixo valor. Um algoritmo ganancioso somaria 8d com três moedas, uma de 6d, uma de 2d e uma de 1d, mas o mesmo podia ser feito com dois *groats*. Esta situação surge quando duplicamos o valor de uma das nossas moedas maiores do que 2d (a de 24d ou a de 4d), porque não existe outra moeda com esse valor. Esta confusão é evitada com as moedas modernas dos EUA, do Reino Unido ou da União Europeia.

Todos os sistemas modernos escolhem de entre um conjunto semelhante de números redondos, 1, 2, 5, 10, 20, 25 e 50 para as suas moedas. Mas será esse o melhor conjunto possível? Os números são relativamente fáceis de somar e combinar, mas serão ideais no que diz respeito a pagar valores baixos usando o mínimo de moedas possível?

Há alguns anos, Jeffrey Shallit, residente em Waterloo, Ontário, levou a cabo uma investigação informática do número médio de moedas que eram necessárias para somar qualquer quantia entre 1 e 100 centimos no sistema americano com diferentes escolhas de moedas. Se usássemos o conjunto atual de moedas de 1, 5, 10 e 25 centimos americanos, o número médio de moedas necessárias para somar qualquer valor entre 1 e 100 centimos seria 4,7. Se só existissem moedas de 1 centimo, seriam necessárias 99 moedas para somar 99 centimos e o número médio necessário para todos os valores do intervalo entre 1 e 100 centimos seria 49,5 – o pior possível. Se existissem apenas moedas de 1 e 10 centimos, essa

média caía para 9. A questão interessante é se conseguimos baixar a média usando um conjunto de quatro moedas com valores diferentes das atuais, 1, 5, 10 e 25. A resposta é "sim" e há dois conjuntos de moedas que poderiam produzir melhores resultados. Se só houvesse moedas de 1, 5, 18 e 29 cêntimos ou 1, 5, 18 e 25 cêntimos, o número médio de moedas necessárias para perfazer qualquer valor entre 1 e 100 cêntimos seria apenas 3,89. A melhor solução usaria o conjunto de moedas de 1, 5, 18 e 25 cêntimos e só é necessária uma pequena mudança – substituir a atual moeda de 10 cêntimos por uma moeda de 18 cêntimos!

Vejamos o que acontece ao sistema britânico ou ao europeu quando os sujeitamos a uma análise semelhante e nos perguntamos que moeda podia ser introduzida nos nossos sistemas para os tornar o mais eficientes possível. Estes dois sistemas têm atualmente moedas de 1, 2, 5, 10, 20, 50, 100 e 200 (ou seja, no Reino Unido temos moedas de 1 e 2 libras em vez de notas de 1 libra). Suponhamos que vamos calcular o número médio de moedas necessárias para somar todos os valores de 1 (*pence* ou cêntimo) a 500 nestes sistemas. O número de moedas necessárias é 4,6. No entanto, este número poderia ser reduzido para 3,92 se acrescentássemos uma moeda com o valor de 133 ou 137 *pence* ou cêntimos.

39 | DECOMPOSIÇÃO DA LEI DAS MÉDIAS

Há três tipos de mentiras – mentiras, mentiras vergonhosas e estatísticas.

Benjamin Disraeli

A nossa análise da geração de energia eólica no Capítulo 35 esconde uma interessante subtilidade acerca da estatística que devemos ter sempre em conta. Todos sabemos que existem três tipos de mentiras – mentiras, mentiras vergonhosas e estatísticas, como disse Disraeli –, mas uma coisa é sabermos que devemos ter em mente um determinado facto, outra é saber exactamente qual é o risco que corremos se não o fizermos. Vejamos como podemos ser enganados por um uso traiçoeiro da estatística no caso da energia eólica. Sabemos que a energia presente no vento é proporcional ao cubo da velocidade do vento, V^3 . Para evitar escrever todas as outras quantidades envolvidas, como a densidade do ar, que não muda, digamos que a energia gerada por unidade de tempo é igual a esta, pelo que $P = V^3$. Por uma questão de simplicidade, estimámos que a velocidade média do vento ao longo de um ano era de 5 metros por segundo, pelo que a quantidade de energia gerada ao longo de um ano seria $5^3 \times 1 \text{ ano} = 125$ unidades.

Na prática, a velocidade média do vento não corresponde sempre à velocidade do vento. Suponhamos que incluímos uma variação muito simples – que a velocidade do vento é zero em metade dos dias do ano e igual a 10 metros por segundo na outra metade. A velocidade média do vento continua a ser $\frac{1}{2} \times 10 = 5$ metros por segundo, mas agora qual é a quantidade de energia gerada? Durante metade do ano deverá ser igual a zero, porque a velocidade do vento é zero; durante a outra metade será $1/0 \times 10 = 50$. O total gerado ao longo do ano é assim de 500 unidades,

muito mais do que obtemos se só considerarmos a velocidade média. Os dias com maior velocidade do vento compensam grandemente em produção energética os dias em que não existe vento. Na realidade, a distribuição das velocidades do vento ao longo de um ano é mais complicada do que aquela aqui apresentada, mas tem a mesma propriedade de produzir muito mais ganho energético quando a velocidade do vento é superior à média do que perda energética quando a velocidade do vento está abaixo da média. Esta situação contraria aquilo que é vulgarmente conhecido como a *lei das médias*. De facto, não se trata de uma lei geral, mas apenas da intuição que muitas pessoas têm de que as variações se equilibram no fim: isto é, que ao longo do tempo haverá tantos ganhos como perdas de dimensões comparáveis acima e abaixo do comportamento médio. Isto só é verdade em relação a variações aleatórias governadas por estatísticas que têm um padrão simétrico especial. Não é o caso do nosso problema de geração de energia eólica e, assim, os ventos com velocidades acima da média têm um efeito muito maior do que os ventos com velocidades abaixo da média. Portanto, tenha cuidado com as médias.

40 | QUANTO TEMPO É PROVÁVEL QUE AS COISAS SOBREVIVAM?

As estatísticas são como um biquíni. O que revelam é sugestivo, mas o que ocultam é vital.

Aaron Levenstein

As estatísticas são ferramentas poderosas. Parecem dizer-nos algo a troco de nada. Frequentemente preveem os resultados das eleições baseando-se apenas nas decisões antecipadas de alguns eleitores num ou dois círculos eleitorais. À primeira vista, parece que estão a tirar as suas conclusões a partir de evidências inexistentes. Um dos meus exemplos preferidos deste aparente truque de magia está relacionado com a previsão do futuro. Se temos uma instituição ou tradição que já existe ao longo de um período determinado, quanto tempo podemos esperar que sobreviva no futuro? A ideia-chave que está na base deste raciocínio é bastante simples. Se vissemos algo num momento determinado aleatoriamente, haveria 95% de probabilidades de estarmos a vê-lo durante os 95% centrais do seu tempo de existência total. Consideremos um período de tempo igual a uma unidade de tempo e tiremos o intervalo médio de 0,95 de modo a que haja um intervalo de 0,025 no início e no fim:

$$0 \longleftarrow 0,025 \longrightarrow A \longleftarrow 0,95 \longrightarrow B \longleftarrow 0,025 \longrightarrow -1$$

Se a situação em questão tiver sido observada no ponto A, o futuro ocupará $0,95 + 0,025 = 0,975$ e o passado ocupará 0,025. Desta forma, o futuro será $975/25 = 39$ vezes mais longo do que o passado. Da mesma forma, se a observar no ponto B, o futuro será apenas uma fração do passado: $1/39$.

Quando observamos algo em algum ponto da história, como a existência da Grande Muralha da China ou da Universidade de Cambridge, se partirmos do pressuposto de que não estamos a observá-lo num momento especial da História, na medida em que não existe um motivo pelo qual devamos estar especialmente perto do seu início ou fim, podemos prever quanto tempo devemos esperar que exista no futuro, com um nível de confiança de 95%. O nosso diagrama mostra-nos os números que precisamos de conhecer. Se uma instituição existir há aproximadamente Y anos e estivermos a observá-la num momento aleatório da História, podemos ter 95% de certeza de que existirá durante pelo menos $Y/39$ anos, mas não mais do que $39 \times Y$ anos¹⁰. Em 2008, a Universidade de Cambridge já existia há 800 anos. Esta fórmula prevê que a sua longevidade futura para além daquela data terá 95% de probabilidades de ser de pelo menos $800/39 = 20,5$ anos, aproximadamente, e no máximo de $800 \times 39 = 31\,200$ anos. A independência dos Estados Unidos foi declarada em 1776. Este país tem 95% de probabilidades de se manter uma nação independente durante mais 5,7 anos, mas menos de 8736 anos. A Humanidade existe há cerca de 250 000 anos, pelo que sobreviverá entre 6410 e 9 750 000 anos com um nível de confiança de 95%. O Universo está a expandir-se há 13,7 mil milhões de anos. Se esta é a duração da sua existência, há 95% de probabilidades de durar outros 351 milhões de anos e menos de 534,3 mil milhões de anos.

Experimente aplicar esta fórmula a coisas que existam há menos tempo. A minha casa foi construída há 39 anos, pelo que posso ter 95% de certeza de que não vai sucumbir a um qualquer desastre aleatório no próximo ano, mas eu (ou qualquer outra pessoa) deveria ficar surpreendido, com 95% de probabilidades, se ela sobreviver durante mais 1521 anos. Pode aplicar esta fórmula a dirigentes de clubes desportivos, empresas, países, partidos políticos, modas e peças de teatro em cena. Também é interessante para fazer previsões. No momento em que estou a escrever este livro, Gordon

¹⁰ É tentador continuar a aplicar esta regra a tudo, mas é preciso ter cuidado. Algumas coisas têm uma esperança média de vida que é determinada por algo mais do que o simples acaso (como a bioquímica). Quando esta regra da probabilidade é aplicada em situações em que um processo não-aleatório determina a duração da mudança, isso pode conduzir a uma conclusão errada acerca do tempo máximo (ou mínimo) de sobrevivência previsto. A fórmula diz que há 95% de probabilidades de uma pessoa de 78 anos viver durante mais 2 a 3042 anos. No entanto, espera-se que a probabilidade de essa pessoa viver mais 50 anos seja precisamente zero. A idade de 78 anos não é um momento aleatório na vida de uma pessoa: está muito mais próximo do fim biológico do que do dia em que a pessoa nasceu. Para uma discussão mais detalhada, consulte o website <http://arxiv.org/abs/0806.3538>.

Brown é Primeiro-Ministro do Reino Unido desde o dia 27 de junho de 2007 – há quase exatamente três meses. Fez ontem o seu primeiro discurso na Conferência do Partido Trabalhista, mas podemos afirmar, com um nível de confiança de 95% que estará em funções durante mais dois ou três dias, mas não mais do que 9 anos e $\frac{3}{4}$. Quando o livro tiver sido publicado, o leitor terá a oportunidade de testar esta previsão!

41 | O PRESIDENTE QUE PREFERIA O TRIÂNGULO AO PENTÁGONO

O triângulo parece não exigir nenhuma capacidade especializada para ser tocado.

O Triângulo, *Wikipédia*

James Garfield foi o 20.º Presidente dos Estados Unidos. O leitor provavelmente nada sabe acerca dele, exceto que, no dia 2 de julho de 1881, ao fim de apenas quatro meses em funções, foi baleado, por um membro descontente do público que queria ter sido nomeado para um cargo na administração federal, e que morreu dez semanas depois. Estranhamente, a bala fatal que se alojou no corpo dele nunca foi encontrada, apesar de Alexander Graham Bell ter sido contratado para inventar um detetor de metais para o efeito. Bell conseguiu inventar o dispositivo em questão, mas não foi muito eficaz, principalmente, suspeita-se, por a cama de Garfield na Casa Branca ter uma estrutura metálica – algo que era muito raro na época – e ninguém suspeitou que fosse essa a causa da falha do instrumento. Na realidade, a causa de morte foi um procedimento médico descuidado que lhe perfurou o fígado. Assim, tornou-se o segundo Presidente a ser assassinado e o que permaneceu em funções durante o segundo período mais curto. Apesar do seu fim trágico, o nome dele ficou para a posteridade devido ao seu invulgar contributo para a matemática.

Originalmente, Garfield tencionava ser professor de matemática depois de se licenciar pela Williams College em 1856. Deu aulas de matemática clássica durante algum tempo e fez uma tentativa malsucedida de se tornar reitor, mas o seu patriotismo e as suas opiniões fortes levaram-no a concorrer a um cargo público e foi eleito para o Senado Estatal do Ohio

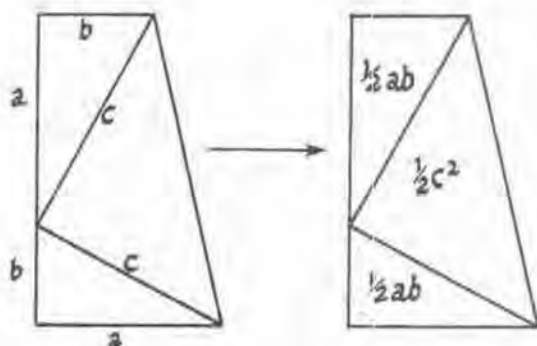
apenas três anos mais tarde, tendo obtido o seu diploma em direito em 1860. Deixou o Senado e juntou-se ao exército em 1861; subiu rapidamente na carreira até se tornar major-general, antes de mudar de rumo e entrar para a Câmara dos Deputados dois anos mais tarde. Ali permaneceu durante dezassete anos até se tornar candidato à Presidência pelo Partido Republicano em 1880, vencendo por pouco o candidato do Partido Democrata, Winfield Hancock, com um programa que prometia melhorar a educação para todos os Americanos. Garfield continua a ser a única pessoa na história a ser eleita Presidente diretamente a partir da Câmara dos Deputados.

O contributo mais interessante de Garfield nada teve que ver com política. Em 1876, enquanto estava em funções, gostava de se reunir com os membros do Congresso para discutir assuntos complexos. Garfield apresentou para entretenimento dos colegas uma nova prova do teorema de Pitágoras para triângulos retângulos. Mais tarde publicou-a no *New England Journal of Education*, onde fez notar que “consideramos ser um assunto em que os membros de ambas as Câmaras estão de acordo, sem distinção partidária”.

Os matemáticos ensinam este teorema aos alunos há mais de 2000 anos e normalmente não se afastam muito das provas dadas por Euclides no seu famoso *Elementos*, uma obra que escreveu em Alexandria em cerca de 300 a.C. Esta prova não foi, de forma alguma, a primeira. Tanto os Babilónios como os Chineses tinham boas provas e os antigos Egípcios estavam familiarizados com o teorema e conseguiam usá-lo nos seus ousados projetos de construção. De todas as provas do teorema de Pitágoras que foram descobertas ao longo dos séculos, a de Garfield é uma das mais simples e a mais fácil de compreender.

Tomemos como exemplo um triângulo retângulo e designemos os seus lados por a , b e c . Façamos uma cópia dele coloquemos as duas cópias lado a lado de modo a formarem uma fenda em forma de V com uma base direita. Agora, unamos as duas pontas mais altas dos dois triângulos, de modo a formar um triângulo assimétrico – uma forma a que chamaremos trapézio (ou, na América, trapezoide). Temos uma figura com quatro lados, com dois dos lados paralelos, conforme se pode ver na figura.

A figura proposta por Garfield é constituída por três triângulos – as duas cópias com que começámos e o terceiro que foi criado quando desenhámos



uma linha a unir os dois pontos mais altos dos dois triângulos. Ele pede-nos simplesmente que calculemos a área do trapézio de duas formas. Primeiro, como um todo, é a altura, $a+b$, vezes a largura média $\frac{1}{2}(a+b)$; pelo que a área do trapézio é $\frac{1}{2}(a+b)^2$. Para se convencer disto, pode mudar a forma do trapézio, transformando-o num retângulo, tornando as duas larguras iguais. A nova largura seria $\frac{1}{2}(a+b)$.

Agora calculemos a área de outra forma. É simplesmente a soma das áreas dos três triângulos retângulos que a constituem. A área do triângulo retângulo é apenas metade da área do quadrado que obteríamos se juntássemos as duas cópias do triângulo na diagonal, pelo que é metade da base vezes a altura. Assim, a área total dos três triângulos é $\frac{1}{2}ba + \frac{1}{2}c^2 + \frac{1}{2}ba = ba + \frac{1}{2}c^2$, conforme mostrado.

Tendo em conta que estes dois cálculos da área deverão resultar no mesmo total, temos:

$$\frac{1}{2}(a+b)^2 = ba + \frac{1}{2}c^2$$

Ou seja,

$$\frac{1}{2}(a^2 + b^2 + 2ab) = ba + \frac{1}{2}c^2$$

Assim, multiplicando por dois:

$$a^2 + b^2 = c^2$$

Tal como Pitágoras disse.

Devíamos pedir a todos os candidatos às eleições americanas que demonstrassem esta prova durante os debates televisivos.

42 | CÓDIGOS SECRETOS NO SEU BOLSO

E ninguém poderá comprar nem vender se não tiver a marca.

Livro do Apocalipse

Os códigos estão sempre relacionados com espões, fórmulas secretas e países em guerra, certo? Errado. Os códigos estão por toda a parte, nos cartões de crédito, nos cheques, nas notas e até mesmo na capa deste livro. Às vezes os códigos desempenham o seu papel tradicional de encriptação de mensagens para os bisbilhoteiros não conseguirem lê-las facilmente ou para impedir terceiros de atacarem a sua conta bancária através da Internet, mas também têm outros usos. As bases de dados têm de ser mantidas livres de corrupção acidental para além de invasões maliciosas. Se alguém inserisse o número do seu cartão de crédito na máquina de pagamento, mas se enganasse num dígito (normalmente acontece trocando dígitos adjacentes, como 43 por 34, ou enganando-se a escrever pares do mesmo dígito, transformando 899 em 889), a sua compra poderia ser cobrada a outra pessoa. Se introduzisse erradamente um número de contribuinte, ou o código de um bilhete de avião, ou um número de passaporte, o erro poderia espalhar-se por parte do mundo eletrónico, criando uma confusão crescente.

O mundo comercial tentou combater este problema criando meios pelos quais estes números importantes podem verificar-se a si próprios, dizendo aos computadores se o número que está a ser introduzido pode ou não ser um número válido para um bilhete de avião ou o número de série válido de uma nota. Há uma variedade de sistemas semelhantes em funcionamento para verificar a validade dos números dos cartões de crédito. A maioria das empresas usa um sistema desenvolvido pela IBM

para cartões de crédito com números de 12 ou 16 dígitos. O processo é um pouco laborioso se for feito manualmente, mas pode ser verificado instantaneamente por uma máquina, que rejeitará o número introduzido se os dígitos que o compõem falharem o teste devido a um erro de introdução ou à má qualidade da falsificação.

Consideremos o número de um cartão Visa imaginário:

4000 1234 5678 9314

Em primeiro lugar, vamos dobrar o valor de cada dígito alternado, avançando da esquerda para a direita, a começar pelo primeiro (ou seja, as posições ímpares), obtendo os números 8, 0, 2, 6, 10, 14, 18, 2. Nos casos em que o número obtido seja composto por dois dígitos (como 10, 14, 18), somamos os dois dígitos (obtendo os números 1, 5, 9), o que tem o mesmo efeito que subtrair 9 a cada um deles. A lista de números é agora 8, 0, 2, 6, 1, 5, 9, 2. Em seguida, somamos todos estes números e depois somamos-lhes os números alternados que não considerámos (os que estavam nas posições pares) da primeira vez (0, 0, 2, 4, 6, 8, 3, 4). A soma resultante, por ordem, será:

$$8+0+0+0+2+2+6+4+1+6+5+8+9+3+2+4=60$$

Para o número do cartão poder se considerado válido, este total tem de ser exatamente divisível por 10, que é o caso deste exemplo. Se o número introduzido tivesse sido 4000 1234 5678 9010, o mesmo cálculo teria gerado o número 53 (porque o número do cartão só difere no antepenúltimo e no último dígitos) e este número não é exatamente divisível por 10. O mesmo procedimento é válido para a verificação da autenticidade da maioria dos cartões de crédito.

Este sistema de verificação deteta muitos erros simples de digitação e de leitura. Deteta todos os erros de introdução de um único dígito e a maior parte das trocas de numerais adjacentes (embora a troca de 90 por 09 passasse despercebida).

Outro código de verificação que encontramos constantemente (e ignoramos completamente, a não ser que sejamos empregados de caixa de um supermercado) é o Código Universal de Produto, ou CUP, que começou

a ser usado em artigos de mercearia em 1973, mas que desde então passou a ser usado em praticamente todos os produtos que encontramos nas lojas. Este é um número com 12 dígitos que é representado por uma fila de barras verticais que um leitor a laser consegue ler facilmente. O CUP é constituído por quatro partes: abaixo das barras, duas sequências distintas de 5 dígitos são dispostas entre dois dígitos isolados. Na caixa da máquina fotográfica digital que tenho em cima da secretária neste momento, por exemplo, lê-se:

0 74101 40140 0

O primeiro dígito identifica o tipo de produto. Os dígitos 0, 1, 6, 7, 9 são usados para todos os tipos de produtos, o dígito 2 é reservado para produtos como queijo, fruta e vegetais que são vendidos a peso, o dígito 3 é usado para medicamentos e produtos relacionados com a saúde, o dígito 4 é para produtos que vão ser sujeitos a um desconto ou que estão relacionados com os cartões de fidelização da loja, e o dígito 5 destina-se a descontos de cupões ou ofertas semelhantes. O bloco de cinco dígitos que se segue identifica o fabricante – que no caso da minha máquina é a Fuji – e os outros cinco são usados pelo fabricante para identificar o produto pelo seu tamanho, cor e outras características para além do preço. O último dígito – neste caso, zero – é o número de controlo. Em alguns casos não é impresso, mas está representado nas barras do código, para o *scanner* poder aceitar ou rejeitar o CUP. O CUP é gerado somando os dígitos nas posições ímpares ($0+4+0+4+1+0=9$), multiplicando esse número por 3 ($3 \times 9 = 27$) e somando ao resultado os números que estavam nas posições pares ($27+7+1+1+0+4+0=40 = 4 \times 10$) e verificando se o número obtido é divisível por 10 – que é claramente o caso.

Restam apenas as barras. O espaço entre os dois dígitos exteriores (os nossos dois zeros) está dividido em sete regiões, que são preenchidas com uma espessura de tinta preta que depende do número que está a ser representado, criando linhas brancas e pretas alternadas. Em cada extremidade do CUP existem duas barras paralelas, ou “guardas”, de espessura igual, que definem a espessura e a escala de separação que é usada para as linhas e para os espaços que as separam. Há um grupo semelhante de quatro barras ao centro, duas das quais por vezes se prolongam abaixo das

outras, e que separam a identificação do fabricante das informações relativas ao produto e que não contêm nenhuma outra informação. As posições e espessura das barras formam um código binário composto por zeros e uns. Um número ímpar de dígitos binários é usado para codificar os dados do fabricante, ao passo que um número par é usado para as informações relativas ao produto. Esta distinção impede que haja confusão entre os dois e permite que um *scanner* consiga ler estes números tanto da direita para a esquerda como da esquerda para a direita, sabendo sempre qual dos blocos está a ler. E o leitor a pensar que a vida era simples.

43 | TENHO UMA PÉSSIMA MEMÓRIA PARA NOMES

O "t" é mudo, como em Harlow.

Margot Asquith, sobre o seu nome ser mal pronunciado por Jean Harlow

Se alguma vez teve de tomar nota do nome de alguém por telefone, sabe como é difícil ter a certeza de como se escreve. Normalmente pedimos à pessoa que solete o seu nome. Lembro-me que o orientador da minha tese de doutoramento, Dennis Sciama, cujo invulgar sobrenome era pronunciado como "Sharma", passava grande parte do tempo a soletrar o seu sobrenome ao telefone para pessoas que não o conheciam.

Há ocasiões em que as mensagens orais ou escritas não podem ser repetidas ou foram escritas erradamente, e queremos minimizar a possibilidade de nos enganarmos no nome de uma pessoa quando o procuramos num arquivo. O sistema mais antigo em uso para minimizar este problema chama-se sistema fonético Soundex e data de 1918, quando foi inventado por dois Americanos, Robert Russel e Margaret Odell, embora desde essa altura tenha sofrido várias pequenas modificações. Foi originalmente concebido para ajudar a garantir a integridade dos dados dos censos que eram recolhidos oralmente e depois passou a ser usado pelas companhias aéreas, pela polícia e pelos sistemas de reserva de bilhetes.

A ideia era codificar os nomes de modo a que pequenas variações da grafia, como Smith e Smyth, ou Ericson e Erickson, que eram foneticamente iguais, fossem codificadas da mesma forma, de modo a que ao introduzir um dos nomes do grupo todas as variantes aparecessem também, garantindo que nenhum dos nomes ficava excluído do sistema de arquivo. Qualquer pessoa que esteja a pesquisar nomes de familiares ou antepassados, especialmente imigrantes com nomes que possam ter sofrido

ligeiras modificações, considerará esta codificação útil. Este sistema procura automaticamente muitas das variantes aproximadas que teríamos, de outra forma, de procurar individualmente, e até variantes de que não nos tínhamos lembrado. Eis a forma como funciona, no caso dos nomes.

1. Manter a primeira letra do nome, qualquer que seja.
2. No resto do nome, eliminar todas as ocorrências das seguintes letras: a, e, i, o, u, h, y, w.
3. Atribuir números a todas as letras restantes, de modo a que:
b, f, p, v passem a ser 1
c, g, j, k, q, s, x, z passem a ser 2
d e t passem a ser 3 e l passe a ser 4
m e n passem a ser 5 e r passe a ser 6
4. Se duas ou mais letras com o mesmo número surgirem seguidas no nome completo original, manter apenas a primeira.
5. Finalmente, registar apenas os primeiros quatro caracteres do resultado obtido. Se o resultado tiver menos de quatro caracteres, adicionar zeros no fim para criar uma sequência com quatro caracteres.

O meu nome é John e este processo transformá-lo-ia em Jn (após os passos 1 e 2), depois em J5 (após o passo 3) e o resultado final seria J500. Se o nome do leitor for Jon, o resultado será exatamente o mesmo. Smith e Smyth passam ambos a ser S530. Ericson, Erickson, Eriksen e Erikson resultam todos no registo E6225.

44 | O CÁLCULO PROLONGA A VIDA

Como professor de matemática, compreendo a importância de mostrar aos alunos a forma como a matemática está relacionada com a vida. A ciência forense proporciona-me uma forma única e inovadora de o fazer. Afinal de contas, o que pode ser mais universal na vida do que a morte? Quando os meus alunos aprendem tudo acerca das velocidades de decomposição e da teoria do embalsamamento, regressam ao estudo do cálculo com um vigor renovado.

Professor Sweeney Todman, Matemático

A diferença entre um amador e um profissional é que o amador tem a liberdade de estudar as coisas que lhe agradam, ao passo que um profissional também tem de estudar o que não lhe agrada. Consequentemente, há partes da formação em matemática que parecerão laboriosas a um aluno, tal como as horas passadas a correr ao frio e à chuva, nos meses de inverno, parecerão pouco apelativas, mas essenciais a um aspirante a atleta olímpico. Quando os meus alunos me perguntavam por que é que tinham de estudar algumas das partes mais complexas e menos cativantes do cálculo, costumava contar-lhes a história que se segue e que foi contada pelo físico russo George Gamow na sua autobiografia, *My World Line*. Esta história fala da experiência notável de um dos amigos de Gamow, um jovem físico de Vladivostok chamado Igor Tamm, que viria a ser um dos distinguidos com o prémio Nobel da física em 1958 pelo seu contributo para a descoberta e compreensão daquilo que é hoje conhecido como o “efeito Cerenkov”.

No período da Revolução Russa, Tamm era um jovem professor de física na Universidade de Odessa, na Ucrânia. A comida era escassa na cidade, por isso deslocou-se a uma aldeia vizinha que aparentemente estava sob o controlo dos comunistas, numa tentativa de trocar algumas colheitas

de prata por algo comestível, como galinhas. Subitamente, a aldeia foi capturada por uma milícia anti-comunista, armada com espingardas e explosivos. Os milicianos desconfiaram de Tamm, que vestia roupas citadinas, e levaram-no à presença do seu líder que exigiu saber quem ele era e o que fazia. Tamm tentou explicar que era um simples professor universitário que tinha ido ali à procura de comida.

"Que tipo de professor?", perguntou o líder da milícia.

"Ensino matemática", respondeu Tamm.

"Matemática?", perguntou o miliciano. "Muito bem! Então dá-me uma estimativa do erro que se obtém ao cortar a série Maclaurin no enésimo termo¹¹. Se responderes corretamente serás libertado. Se errares, levas um tiro!"

Tamm ficou estupefacto. Com uma arma apontada à cabeça, um pouco nervoso, conseguiu encontrar a resposta para o problema – um aspeto complicado da matemática que os alunos aprendem no primeiro ano do curso de cálculo numa licenciatura em matemática. Mostrou a resposta ao líder da milícia, que a analisou e declarou: "Correto! Podes voltar para casa!"

Tamm nunca descobriu quem era aquele estranho miliciano. Provavelmente, acabou a trabalhar como diretor de controlo de qualidade do ensino universitário.

Há muito tempo, dois aviadores arranjarão asas. Dédalo voou em segurança a uma altura média e foi devidamente honrado na aterragem. Ícaro voou em direção ao sol, até que a cera que lhe unia as asas derreteu e o seu voo terminou num fiasco. As autoridades clássicas dizem-nos, claro, que ele estava apenas "a fazer acrobacias"; mas eu prefiro pensar nele como o homem que evidenciou um grave defeito de construção das máquinas voadoras do seu tempo.

Arthur S. Eddington

Há muitas coisas que se deslocam com um movimento semelhante ao bater de asas: os pássaros e borboletas com as asas, as baleias e rubarões com as caudas, os restantes peixes com as barbatanas. Em todas estas situações há três fatores importantes que determinam a facilidade e eficácia do movimento. Primeiro, há o tamanho – as criaturas maiores são mais fortes e podem ter asas e barbatanas maiores, que atuam sobre maiores volumes de ar e água. A seguir há a velocidade – aquela a que conseguem voar ou nadar diz-nos a velocidade com que estão a interagir com o meio em que se deslocam e a resistência que este cria, fazendo-as abrandar. Em terceiro lugar, há o ritmo a que conseguem agitar as asas ou barbatanas. Existirá um fator comum que nos permita considerar todos os diferentes movimentos dos pássaros e peixes sob o mesmo ângulo?

Como provavelmente já adivinhou, existe, de facto, esse fator. Quando os cientistas e os matemáticos se deparam com uma diversidade de exemplos de um fenómeno, como o voo ou a natação (exemplos esses que diferem nos pormenores, mas que retêm uma semelhança básica), frequentemente tentam classificar os diferentes exemplos avaliando uma quantidade que constitua um número puro. Com isto, quero dizer que não tem unidades, como acontece com as massas ou velocidades (extensões por unidade de tempo). Isto garante que se mantém igual se as unidades

usadas para medir as quantidades mudarem. Desta forma, uma vez que o valor numérico de uma distância percorrida muda de 10 000 para $6\frac{1}{4}$, quando se altera a unidade de medida de metros para milhas, a proporção entre duas distâncias – como a distância percorrida a dividir pelo comprimento do passo – não mudará se medir a distância e o comprimento do passo *na mesma unidade*, porque se trata simplesmente do número de passos necessários para percorrer a distância em questão.

No nosso caso há uma forma de combinar os três fatores críticos – a frequência de batidas por unidade de tempo, f , o tamanho das batidas, T , e a velocidade de deslocação, V – de modo a obter uma quantidade que seja um número puro*. Esta combinação traduz-se por fT/V e é designada por “número de Strouhal”, em homenagem a Vincenc Strouhal (1850-1922), um físico Checo da Charles University, em Praga.

Em 2003, Graham Taylor, Robert Nudds e Adrian Thomas, da Universidade de Oxford, demonstraram que se avaliarmos o valor deste número Strouhal $St = fT/V$ aplicado a muitas variedades de animais a voar e a nadar às suas velocidades de deslocação habituais (em vez de breves impulsos na perseguição de presas ou a fugir de um ataque), verificamos que se enquadram numa gama relativamente reduzida de valores que se pode considerar que caracterizam os resultados das diferentes histórias evolucionárias que levaram à existência destes animais. Eles consideraram um grande número de animais, mas para o nosso exemplo vamos selecionar apenas alguns animais diferentes para perceber como esta unidade se manifesta face à diversidade superficial.

No caso de um pássaro em voo, f será a frequência das batidas de asas por segundo, T será a amplitude global das duas asas e V será a velocidade de deslocação. Um falcão típico tem um valor f de cerca de 5,6 batidas por segundo, uma amplitude de asas de cerca de 0,34 metros e uma velocidade de deslocação de cerca de 8 metros por segundo, o que resulta em $St(\text{falcão}) = (5,6 \times 0,34)/8 = 0,24$. Um morcego comum tem um valor $V = 6$ metros por segundo, uma amplitude de asas de 0,26 metros e uma taxa de batidas de asa de 8 vezes por segundo, o que resulta num número Strouhal de $St(\text{morcego}) = (8 \times 0,26)/6 = 0,35$. Se repetirmos o cálculo em relação

* A unidade de frequência f é 1/tempo, a de dimensão, L , é o comprimento e a de velocidade, V , é comprimento/tempo, pelo que a combinação fL/V não tem unidades: é um número puro sem dimensão.

a quarenta e duas aves, morcegos e insetos diferentes, o valor St que alcançamos inclui-se sempre no intervalo 0,2-0,4. O resultado que obtiveram com o estudo das espécies marinhas levou à mesma conclusão. Outros estudos mais extensivos foram levados a cabo por Jim Rohr e Frank Fish em San Diego e West Chester, na Pensilvânia, para investigar esta quantidade em relação aos peixes, tubarões, golfinhos e baleias. Descobriu-se que a maioria (44%) enquadrava-se no intervalo entre 0,23 e 0,28, mas que o valor global ia de 0,2 a 0,4, tal como acontecera no estudo dos animais voadores.

Este estudo também pode ser aplicado aos seres humanos. Um bom nadador do sexo masculino nadará 100 metros em 60 segundos, portanto $V = 5/3$ metros por segundo, e usa cerca de 54 braçadas completas dadas com cada braço (pelo que a frequência das braçadas será de 0,9 por segundo), com um alcance de braços debaixo de água de cerca de 0,7 metros. O resultado é $St(\text{nadador humano}) = (0,9 \times 2/3) / 5/3 = 0,36$, o que nos aproxima mais das aves e dos peixes do que imaginávamos. No entanto, a melhor nadadora do mundo é sem dúvida a nadadora australiana de longa distância, Shelley Taylor-Smith, que ganhou os campeonatos mundiais de natação sete vezes. Percorreu a nado 70 km no mar em 20 horas, com uma frequência média de 88 braçadas por minuto. As suas braçadas têm um alcance de um metro, o que lhe dá o incrível número de Strouhal de 1,5 – é quase uma sereia.

46 | O SEU NÚMERO FOI ESCOLHIDO

Introduza o seu código postal para encontrar as lojas de artigos carnavalescos perto de si.

Loja de artigos carnavalescos no Reino Unido

A vida parece ser constantemente definida por números. Temos de nos lembrar de códigos PIN, números de contas, códigos de acesso e uma quantidade infindável de números de referência para todas as instituições e departamentos governamentais à face da Terra, incluindo alguns de que ninguém nunca ouviu falar. Às vezes pergunto-me se algum dia esgotaremos os números. Um número que nos é familiar e que nos classifica geograficamente (de uma forma aproximada) é o código postal. O meu é CB3 9LN e, juntamente com o número da porta da minha casa é suficiente para que todo o meu correio seja entregue sem falhas, embora insistamos em acrescentar nomes de ruas e de cidades como informação complementar, talvez por nos parecer mais humano. O meu código postal segue um padrão quase universal no Reino Unido: é composto por quatro letras e dois números. O posicionamento das letras e dos números não é importante, embora na prática o seja porque as letras também designam os centros regionais de tratamento e distribuição do correio (CB corresponde a Cambridge). Mas não nos detenhamos neste pormenor – de certeza que os correios não se importariam com ele se se esgotassem os códigos postais com seis dígitos – e perguntemos simplesmente quantos códigos postais com esta forma podem existir. Há 26 opções, de A a Z, para cada um dos espaços designados para as letras e 10 escolhas, de 0 a 9, para cada um dos numerais. Se estes são escolhidos independentemente, o número total de códigos postais diferentes que seguem o padrão atual

é igual a $26 \times 26 \times 10 \times 10 \times 26 \times 26$, que é igual a 45 697 600, quase 46 milhões. Atualmente, estima-se que o número de casas do Reino Unido seja cerca de 26 222 000, ou seja, pouco mais de 26 milhões, e prevê-se que aumente para cerca de 28,5 milhões até ao ano 2020. Portanto, até mesmo os nossos códigos postais relativamente pequenos têm capacidade mais do que suficiente para lidar com o número de casas e de dar um número de identificação único a cada uma se vier a ser necessário.

Se quisermos atribuir uma identificação única a cada indivíduo, a fórmula dos códigos postais deixa de ser suficiente. Em 2006, estimava-se que a população do Reino Unido era de 60 587 000, cerca de 60,5 milhões de pessoas, um número muito superior ao de códigos postais. A coisa mais parecida que temos como um número de identificação é o número de Segurança Social, que é usado por vários departamentos governamentais para nos identificar – a possibilidade de todos estes departamentos cruzarem os seus dados usando este número é a possibilidade que mais alarma muitos grupos que defendem as liberdades civis. O número de Segurança Social no Reino Unido segue um padrão com a forma NA 123456 Z, composto por seis numerais e três letras. Tal como no exemplo anterior, podemos calcular facilmente a quantidade de números de Segurança Social que esta fórmula permite:

$$26 \times 26 \times 10 \times 10 \times 10 \times 10 \times 10 \times 10 \times 26$$

Este é um número *grande* – 17 576 000 000 ou, por extenso, dezassete mil quinhentos e setenta e seis milhões, um número muito superior ao da população do Reino Unido (e também superior ao número projetado para 2050, que é de 75 milhões). De facto, a população do mundo inteiro ronda atualmente os 6,65 mil milhões e estima-se que venha a atingir os 9 mil milhões no anos de 2050. Como veem, temos números – e letras – de sobra.

O valor dos seus investimentos tanto pode baixar como subir.

Aconselhamento financeiro aos consumidores
no Reino Unido

Recentemente, o leitor terá descoberto que o valor dos seus investimentos tanto pode subir como descer. Suponhamos que quer jogar pelo seguro e investir o seu dinheiro numa conta poupança normal, com uma taxa de juro fixa ou em mudança lenta. Quanto tempo será necessário para duplicar o seu investimento? Embora não haja nada neste mundo tão certo como a morte ou os impostos, esqueçamos isso por agora e optemos por uma regra simples e prática para calcular o tempo necessário para duplicar o seu investimento.

Começemos por investir um montante M numa conta poupança com uma taxa de juro anual j (pelo que 5% de juros correspondem a $j = 0,05$). Esta terá aumentado para $M \times (1+j)$ ao fim de 1 ano, para $M \times (1+j)^2$ ao fim de 2 anos, para $M \times (1+j)^3$ ao fim de 3 anos, etc. Ao fim de n anos, as suas poupanças terão atingido um montante igual a $M \times (1+j)^n$. Este será igual ao dobro do seu investimento inicial, ou seja, $2M$, em que $(1+j)^n = 2$. Se observarmos os logaritmos naturais desta fórmula e observarmos que $\ln(2) = 0,69$ aproximadamente, e que $\ln(1+j)$ é aproximadamente igual a j quando j é muito menor do que 1 (o que é sempre o caso – atualmente, j é mais ou menos igual a 0,05 ou 0,06 no Reino Unido), o número de anos necessários para duplicar o seu investimento é determinado pela fórmula simples $n = 0,69/j$. Arredondemos 0,69 para 0,7 e consideremos j como sendo J por cento, de modo a que $J = 100j$, obtemos uma regra prática segundo a qual¹²:

$$n = 70/R$$

Isto mostra, por exemplo, que quando a taxa J é de 7%, precisamos de cerca de 10 anos para duplicar o nosso investimento, mas que bastará as taxas de juro caírem para 3,5% para serem necessários 20 anos.

No momento seguinte, Alice tinha atravessado o espelho e saltado com ligeireza para a sala do espelho.

Lewis Carroll

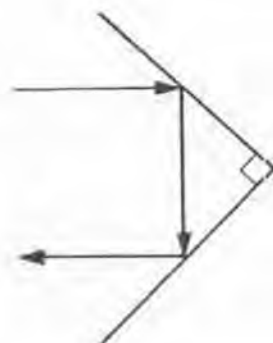
Nós nunca vemos o nosso próprio rosto, exceto quando olhamos para o espelho. A imagem que vemos será verdadeira? Basta uma experiência simples para descobrir. Quando o espelho da casa de banho ficar coberto de vapor, desenhe um círculo em volta da imagem do seu rosto no vidro. Meça o diâmetro desse círculo com a distância entre o seu indicador e o seu polegar, e compare-o com o tamanho real do seu rosto. Descobrirá sempre que a imagem no espelho tem metade do tamanho real do seu rosto. Independentemente da distância a que se coloque do espelho, a imagem do espelho corresponderá sempre a metade da realidade.

É muito estranho. Habituíamo-nos de tal forma à imagem que vemos no espelho enquanto fazemos a barba ou escovamos o cabelo quase todos os dias da nossa vida que nos tornámos imunes à grande diferença entre a realidade e a realidade percebida. Não há nada de misterioso na ótica desta situação. Quando olhamos para um espelho plano, forma-se sempre uma imagem "virtual" do nosso rosto à mesma distância atrás do espelho que estamos a contemplar. Desta forma, o espelho fica sempre a metade da distância entre o rosto real e a imagem virtual. Obviamente, a luz não atravessa o espelho para criar uma imagem do outro lado; a imagem apenas *parece* estar a vir de trás do espelho. Se avançar na direção de um espelho plano, constatará que a sua imagem parece estar a avançar na sua direção ao dobro da velocidade a que está a caminhar.

Outra coisa estranha acerca da imagem no espelho é o facto de ter mudado de orientação. Se segurar a escova de dentes com a mão direita, ela aparecerá na sua mão esquerda na imagem refletida no espelho.

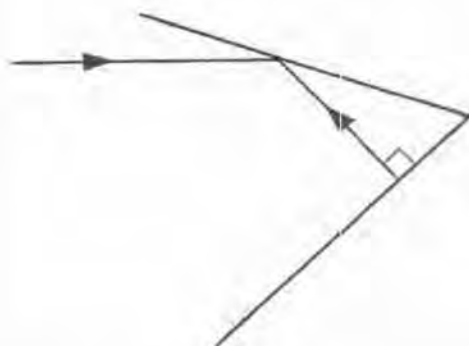
A imagem é invertida da esquerda para a direita, mas não é invertida de cima para baixo: se estiver a observar o seu reflexo num espelho de mão e virar o espelho 90 graus no sentido dos ponteiros do relógio, a sua imagem não se alterará.

Se erguer em frente ao espelho uma folha transparente com algo escrito, acontecerá algo diferente. A escrita não é invertida pelo espelho se tivermos a transparência voltada para nós, de modo a permitir-nos ler o que tem escrito. Se erguermos uma folha de papel também virada para nós, não poderemos lê-la porque é opaca. O espelho permite-nos ver o verso de um objeto que não seja transparente apesar de estarmos à frente dele. No entanto, para podermos ver a parte da frente do objeto, teríamos de o virar. Se o virarmos em torno do seu eixo vertical para o voltarmos para o espelho, invertemos os lados direito e esquerdo. A inversão do espelho entre a esquerda e a direita é produzida por esta alteração no objeto. Se virarmos um livro que estejamos a ler em torno do seu eixo horizontal de modo a que fique com as páginas viradas para o espelho, estas vão aparecer refletidas de cabeça para baixo porque na verdade foram invertidas de cima para baixo e não da esquerda para a direita. Quando não estamos na presença de um espelho, não obtemos estes efeitos porque só conseguimos ver a frente do objeto que estamos a observar e depois de o virarmos só conseguimos ver o seu verso. As letras impressas na página do livro estão invertidas da esquerda para a direita porque virámos o livro em torno do seu eixo vertical para o voltarmos para o espelho. As letras não são viradas de cima para baixo, mas podemos produzir esse efeito no reflexo no espelho se virarmos o livro em torno do seu eixo horizontal, de baixo para cima.



A história não acaba aqui. Há outras coisas interessantes que acontecem (como os ilusionistas bem sabem) se tiver dois espelhos planos. Coloque-os em ângulo reto, de modo a formarem um L e olhe para o canto do L. Este efeito pode ser produzido com um espelho de um toucador que tenha espelhos laterais ajustáveis.

Observe a sua imagem, ou a das páginas de um livro neste par de espelhos dispostos em ângulo reto e descobrirá que a imagem *não* se inverte da esquerda para a direita. A escova de dentes aparece na sua mão direita se estiver realmente na sua mão direita. Na realidade é muito confuso tentarmos barbear-nos ou escovar o cabelo em frente a um sistema deste tipo porque o cérebro faz automaticamente a troca entre a esquerda e a direita. Se alterar o ângulo que une os espelhos, reduzindo-o gradualmente para menos de 90 graus, observará que se produz um efeito estranho quando alcança o ângulo de 60 graus. A imagem obtida é exatamente igual à que seria obtida a olhar para um simples espelho plano e é invertida da esquerda para a direita. A inclinação de 60 graus dos espelhos garante que um feixe de luz que incida sobre um espelho regresse exatamente pelo mesmo caminho e crie o mesmo tipo de imagem virtual que veria num espelho plano simples.



Ele é o organizador de metade do que é mau e de praticamente tudo o que não é detetado [em Londres]. É um gênio, um filósofo, um pensador abstrato. Tem um cérebro excepcional.

Sherlock Holmes, in *O Problema Final*

Houve um tempo – e para alguns esse tempo ainda existe – em que o matemático mais famoso entre o público em geral era uma personagem de ficção. O Professor James Moriarty foi uma das personagens secundárias criadas por Arthur Conan Doyle para o Sherlock Holmes. O “Napoleão do Crime” foi um adversário digno de Holmes e até foi preciso por vezes recorrer ao talento de Mycroft Holmes para vencer os seus planos maléficos. Este vilão aparece em apenas duas histórias de Holmes, *O Problema Final* e *O Vale do Terror*, mas está frequentemente presente nos bastidores, como no caso de “A Liga dos Ruivos” em que planeia uma engenhosa artimanha para permitir aos seus cúmplices, liderados por John Clay, escavarem um túnel até ao interior do cofre de um banco a partir da cave de uma loja de penhores perto do banco.

Sabemos um pouco acerca da carreira de Moriarty pela descrição que Holmes faz dele. Diz-nos ele que:

É um homem de boas famílias e que teve uma excelente educação, dotado pela natureza com capacidades matemáticas fenomenais. Aos vinte e um anos de idade, escreveu um tratado sobre o teorema binomial, que esteve em voga na Europa. A qualidade desse texto valeu-lhe o lugar de presidente do departamento de Matemática numa das nossas pequenas universidades e tinha, ao que parecia, uma brilhante carreira pela frente. Mas o homem tinha tendências hereditárias do tipo mais diabólico. Corria-lhe no sangue um impulso criminoso que, em vez de se modificar, aumentou e se tornou infinitamente mais poderoso graças aos seus extraordinários poderes mentais. Surgiram rumores sombrios em torno dele na universidade e viu-se obrigado a abandonar o cargo e a vir para Londres.



Professor James Moriarty

Mais tarde, n'O *Vale do Terror*, Holmes revela um pouco mais acerca da carreira académica de Moriarty e da sua versatilidade. Se, por um lado, os seus primeiros trabalhos foram dedicados a séries matemáticas, vinte e quatro anos mais tarde vemo-lo a dedicar-se ativamente ao estudo avançado da astronomia dinâmica.

Pois não é ele o autor de *A Dinâmica de um Asteroide*, um livro que ascende a tais alturas da matemática pura que se diz que nenhum homem da imprensa científica é capaz de o criticar?

Conan Doyle fez um uso cuidadoso de acontecimentos e locais reais na construção das suas histórias e é possível fazer uma estimativa aproximada da identidade real do vilão em quem o Professor Moriarty se baseou. O principal candidato é Adam Worth (1844-1902), um alemão que passou toda a vida na América e que se especializou em crimes audaciosos e engenhosos. De facto, um detetive da Scotland Yard da época, Robert Anderson, chegou mesmo a chamar-lhe o "Napoleão do mundo do crime". Tendo começado por cometer pequenos furtos, evoluiu e começou a organizar assaltos em Nova Iorque. Foi apanhado e preso, mas rapidamente fugiu da cadeia e regressou à vida de crime, expandindo a sua atividade para incluir assaltos a bancos e a libertação do arrombador de cofres Charley Bullard da prisão de White Plains por meio de um túnel. À semelhança do que aconteceu no conto "A Liga dos Ruivos", em novembro de 1869,

com a ajuda de Bullard, assaltou o Boylston National Bank em Boston, escavando um túnel para o interior do cofre do banco a partir de uma loja vizinha. Para escaparem aos detetives da agência Pinkerton, Worth e Bullard fugiram para Inglaterra e, em breve, começaram a fazer assaltos também aí e em Paris, para onde se mudaram em 1871. Worth comprou várias propriedades magníficas em Londres e estabeleceu uma grande rede criminosa, para garantir que estava sempre a uma distância segura dos seus crimes. Aqueles que trabalhavam para ele nunca chegaram a saber o seu nome (usava frequentemente o pseudónimo Henry Raymond), mas foi-lhes transmitida a noção de que não deveriam fazer uso de nenhum tipo de violência nos crimes que cometessem em seu nome. Worth acabou por ser apanhado quando foi visitar Bullard à prisão e esteve preso durante sete anos em Leuven, na Bélgica, acabando por ser libertado em 1897 por bom comportamento. Procedeu imediatamente ao roubo de joias para financiar o seu regresso à vida normal e, através dos bons ofícios da agência de detetives Pinkerton em Chicago, providenciou a devolução de um quadro, "The Duchess of Devonshire", à galeria Agnew & Sons em Londres a troco de uma "recompensa" de \$25 000. Depois disso, Worth regressou a Londres e lá viveu com a família até à sua morte em 1902. A sua campa encontra-se no cemitério de Highgate com o nome Henry J. Raymond.

De facto, Worth tinha roubado aquele quadro de Georgiana Spencer (uma mulher de grande beleza e, segundo se crê, parente da Princesa Diana pelo lado da família Spencer) pintado por Gainsborough* da galeria de Agnew em Londres em 1876 e levou-o consigo para toda a parte durante muitos anos, em vez de o vender. Este facto é um dos principais indícios de que o Professor James Moriarty e Adam Worth eram a mesma pessoa.

No livro *O Vale do Terror*, Moriarty é entrevistado pela polícia na sua casa. Pendurado na parede está um quadro intitulado "La Jeune fille à l'agneau" – jovem com um cordeiro – um trocadilho que faz referência à galeria Agnew que perdeu o quadro, embora nunca se tenha conseguido provar que foi Worth quem o roubou. Mas infelizmente, tanto quanto sei, Worth nunca escreveu um tratado sobre o teorema binomial nem uma monografia sobre a dinâmica dos asteroides.

* O quadro encontra-se atualmente na National Gallery of Art, em Washington D. C., e pode ser visto online em http://commons.wikimedia.org/wiki/Image:Thomas_Gainsborough_Georgiana_Duchess_of_Devonshire_1783.jpg

50 | MONTANHAS-RUSSAS E SAÍDAS DAS AUTOESTRADAS

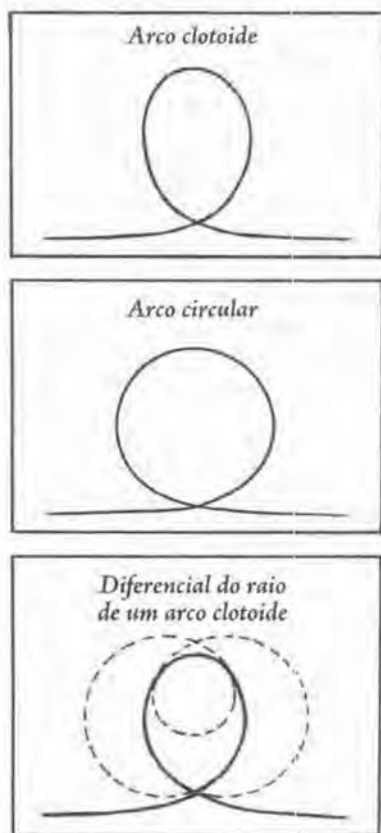
Tudo o que sobe, desce.

Anónimo

Alguma vez andou numa daquelas montanhas-russas assustadoras que descrevem um arco vertical? Poderá ter pensado que o percurso traça um arco circular, mas isso nunca acontece porque se os passageiros chegassem ao topo do arco com uma velocidade suficiente para não caírem dos carros quando atingissem o ponto mais alto (ou, pelo menos, para não ficarem seguros apenas pelos cintos de segurança), a força gravitacional máxima sentida pelos passageiros quando o carro chegasse ao ponto mais baixo seria perigosamente elevada.

Vejam os que acontece se o arco for circular e tiver um raio r e o carro com a carga máxima tiver uma massa m . O carro começará a andar a uma velocidade moderada à altura a (de valor superior a r) acima do chão e efetuará uma descida íngreme até à base do arco. Se ignorarmos todos os efeitos da fricção ou da resistência do ar sobre o movimento do carro, este chegará à base do arco a uma velocidade $V_b = \sqrt{2ga}$. Depois subirá até ao ponto mais elevado do arco. Se alcançar esse ponto a uma velocidade V_t , precisará de uma quantidade de energia igual a $2mgr + \frac{1}{2}mV_t^2$ para vencer a força gravitacional e ascender a uma altura vertical de $2r$ até ao cimo do arco e chegar lá a uma velocidade V_t . Uma vez que a energia total do movimento não pode ser criada nem destruída, temos de ter (a massa do carro, m , cancela todos os termos):

$$gh = \frac{1}{2}V_b^2 = 2gr + \frac{1}{2}V_t^2$$



No topo do arco circular, a força total exercida sobre os passageiros na subida, e que os impede de cair do carro, é a força ascendente do movimento num círculo de raio r menos o seu peso que exerce uma força descendente; ou seja, se a massa do passageiro for M , então:

$$\text{Força ascendente total no topo} = M V_t^2 / r - Mg$$

Este valor tem de ser positivo para impedir que o passageiro caia, pelo que $V_t^2 > gr$.

Se voltarmos a analisar as equações da página 141, concluímos que temos de ter $a > 2,5r$. Isto significa que se for impulsionado apenas com a força da gravidade, tem de começar num ponto pelo menos 2,5 vezes mais alto do que o topo do arco para conseguir chegar ao cimo dele a uma velocidade suficiente para não cair do carro. Mas este é um grande

problema. Se começar num ponto suficientemente elevado para que isto seja possível, chegará ao ponto mais baixo do arco a uma velocidade de $V_b = \sqrt{(2ga)}$, que será superior a $\sqrt{2g(2,5r)} = \sqrt{5gr}$. Quando começa a descrever um arco circular na direção da base, sentirá uma força descendente igual ao seu peso *mais* a força centrífuga, que é igual a:

$$\text{Força descendente total na base} = Mg + M V_b^2 / r > Mg + 5Mg$$

Destá forma, a força descendente total exercida sobre os passageiros na base ultrapassará em seis vezes o seu peso (uma aceleração de 6-g). A maioria dos passageiros, a não ser que sejam astronautas ou pilotos fora de serviço com fatos espaciais, perderão a consciência ao serem sujeitos a esta força. O cérebro deixa de receber oxigénio. Normalmente, os carrosséis destinados a crianças procuram manter as suas acelerações abaixo dos 2-g e aqueles que são destinados aos adultos alcançam acelerações máximas de 4-g.

Segundo este modelo, as montanhas-russas circulares parecem ser uma impossibilidade prática, mas se analisarmos mais atentamente as duas condições – estar sujeito a uma força ascendente suficientemente grande no topo para não cair, mas evitar sofrer forças descendentes fatais quando chega à parte mais baixa – haverá maneira de mudar a forma das montanhas-russas para satisfazer ambas as condições?

Quando nos movemos num círculo com um raio R a uma velocidade V , sentimos uma aceleração centrífuga de V^2/R . Quanto maior for o raio do círculo e, conseqüentemente, menos acentuada a curva, menor será a aceleração sentida. Na montanha-russa, a aceleração V_t^2/r no topo é o que nos impede de cair, sobrepondo-se ao nosso peso Mg que atua em sentido descendente, pelo que é conveniente que seja grande, o que pressupõe um valor r pequeno no topo. Por outro lado, quando estamos no ponto mais baixo, a força circular é o que está a criar os 5-g adicionais de aceleração e podemos reduzir esse valor descrevendo um círculo menos fechado com um raio maior. Isto pode ser conseguido aumentando a altura da montanha-russa e reduzindo a sua largura, de modo a parecer composta por duas partes de dois círculos distintos, tendo o que forma a parte superior um raio menor do que o que forma a parte inferior. A curva habitualmente usada que tem este aspeto chama-se "clotoide", e tem uma curvatura que diminui à medida que avançamos ao longo dela em

proporção com a distância percorrida. Esta curva foi incorporada pela primeira vez na configuração das montanhas-russas em 1976, pelo engenheiro alemão Werner Stengel, na célebre "Revolution" no parque de diversões Six Flags Magic Mountain na Califórnia.

As clotoïdes têm outra característica interessante que levou à sua inclusão na concepção das saídas de algumas autoestradas ou em linhas de caminho de ferro. Se um condutor estiver a circular numa saída com esta curva, desde que mantenha uma velocidade constante, pode simplesmente mover o volante sem alterar a velocidade de rotação. Se a curva tivesse uma forma diferente, teria de ajustar constantemente o movimento do volante ou a velocidade do carro.



Não sei que armas vão ser usadas na III Guerra Mundial, mas a IV Guerra Mundial vai ser feita com paus e pedras.

Albert Einstein

A primeira bomba atômica explodiu no teste de Trinity no Novo México, nos EUA, a 210 milhas a sul de Los Alamos, no dia 16 de julho de 1945. Foi um momento decisivo na história da Humanidade. A criação deste dispositivo deu aos seres humanos a capacidade de destruir toda a vida humana e de produzir consequências fatais a longo prazo. Subsequentemente, deu-se uma corrida às armas que produziu um aumento da produção energética destas bombas quando os EUA e a União Soviética procuraram demonstrar a sua capacidade de produzir explosões cada vez mais devastadoras. Embora só dois dispositivos deste tipo tenham sido usados em conflitos*, as consequências para ecologia e para a saúde humana que esta época de testes provocou na atmosfera, no solo, no subsolo e nos meios subaquáticos ainda se fazem sentir.

Estas explosões foram muito fotografadas na altura e produziam uma bola de fogo característica e uma nuvem de detritos que acabou por simbolizar as consequências da guerra nuclear. A familiar nuvem em forma de cogumelo** forma-se por um motivo. Um grande volume de gases muito quentes com uma baixa densidade forma-se com grande pressão perto do nível do chão – as bombas atômicas e nucleares eram geralmente detonadas acima do nível do chão para maximizar os efeitos da onda explosiva

* A bomba *Little Boy* largada em Hiroshima era um dispositivo de fissão nuclear com 60 kg de urânio-235 e produziu uma explosão equivalente a 13 quilotoneladas de TNT, causando a morte imediata a aproximadamente 80 000 pessoas; a bomba *Fat Man*, largada em Nagasaki, era uma bomba de fissão nuclear com 6,4 kg de plutónio-239, equivalente à explosão de 21 quilotoneladas de TNT. Cerca de 70 000 pessoas tiveram morte imediata.

** Este nome tornou-se popular no início da década de 1950, mas a comparação entre os padrões de dispersão dos detritos e os cogumelos data pelo menos das manchetes dos jornais de 1937.

em todas as direções. Tal como as bolhas que sobem na água a ferver, os gases aceleram em direção ao ar mais denso a maior altitude, criando remoinhos turbulentos que se curvam para baixo nas extremidades, enquanto os detritos adicionais e o fumo sobem na vertical, formando uma coluna ao centro. O material no centro da detonação é vaporizado e aquecido para dezenas de milhões de graus, produzindo uma grande quantidade de raios-X, que colidem com os átomos e moléculas do ar acima dele e os carregam de energia, produzindo um clarão de luz branca cuja duração depende da magnitude da explosão inicial. Quando a parte frontal da coluna sobe, gira como um tornado e puxa os materiais depositados no solo, formando o "caule" da forma do cogumelo que se ergue; a densidade diminui à medida que se espalha e acaba por ficar com a mesma densidade do ar acima de si. Nesse momento, para de subir e dispersa-se para os lados, de modo a que todos os materiais erguidos do chão voltam para trás e descem, criando uma vasta extensão de poeiras radioativas.

As explosões comuns de TNT, ou de outras armas não-nucleares, têm uma aparência bastante diferente por causa das temperaturas mais baixas criadas no início de uma explosão puramente química. Estas resultam numa mistura de gases explosivos em vez da nuvem organizada em forma de cogumelo.

Um dos pioneiros do estudo da forma e das características das grandes explosões foi o notável matemático de Cambridge Geoffrey (G. I.) Taylor. Taylor escreveu um relatório confidencial sobre o carácter esperado da explosão da bomba atômica em junho de 1941. Ficou conhecido por um público mais vasto depois de a revista *Life*, dos EUA, ter publicado uma sequência de fotografias sucessivas do teste de Trinity em 1945 no Novo México. A produção energética desta e de outras explosões americanas da bomba atômica ainda era ultrassecreta, mas Taylor mostrou como apenas algumas linhas de álgebra lhe permitiam (e a qualquer pessoa com algumas noções simples de matemática) calcular a energia aproximada de uma explosão, simplesmente observando as fotografias.

Taylor conseguiu calcular a distância esperada do limiar da explosão em qualquer momento após a detonação, fazendo notar que só dependia de duas coisas: da energia da explosão e da densidade do ar circundante que tem de atravessar. Só há uma apresentação possível para esta situação¹³, que é aproximadamente:

$$\text{Energia da bomba} = \frac{\text{densidade do ar} \times (\text{distância do limiar da explosão})^3}{(\text{tempo desde a detonação})^2}$$

As fotografias publicadas mostravam a explosão em momentos diferentes após a detonação e a hora a que foram tiradas era inscrita ao lado de cada fotografia com uma escala ao longo da parte inferior para determinar o seu tamanho. A partir da primeira fotografia, Taylor observou que ao fim de 0,006 segundos a onda explosiva tinha alcançado um raio de aproximadamente 80 metros. Sabemos que a densidade do ar é de 1,2 kg por metro cúbico, pelo que a equação nos diz que a energia libertada foi de 10^{14} joules, o que é equivalente a 25 000 toneladas de TNT. Em comparação, sabemos que o terramoto que ocorreu na Índia em 2004 libertou uma energia equivalente a 475 milhões de toneladas de TNT.

Conseguimos identificar as pessoas do norte da Europa por caminharem mais depressa do que é adequado para um simples passeio.

Guia de viagens de Benidorm

Caminhe numa rua muito movimentada e reparará que todas as pessoas à sua volta parecem estar a caminhar mais ou menos à mesma velocidade. Há algumas pessoas que caminham um pouco mais apressadas e outras que caminham mais devagar, talvez devido à idade, a algum problema de saúde ou por estarem a usar calçado pouco prático. Quando caminhamos, mantemos sempre parte da superfície do pé em contacto com o chão e esticamos a perna quando levantamos o pé do chão. De facto, as regras da marcha determinam que estas são as características da modalidade e o que a distingue da corrida: o seu não comprimento resulta em avisos e, em última análise, na desqualificação de uma competição. Enquanto caminhamos, levantamos e baixamos as ancas à medida que o nosso centro descreve um ligeiro arco circular de cada vez que damos um passo completo. Assim, se o comprimento da sua perna, do chão até às ancas, for C , estará a produzir uma aceleração igual a v^2/C no sentido ascendente, em direção ao centro desse arco de movimento circular. Este valor não pode ser superior ao da aceleração causada pela gravidade, g , que nos empurra na direção do chão (caso contrário levantaríamos voo!), e assim $g > v^2/C$ o que nos leva a deduzir que a velocidade máxima de uma marcha normal é aproximadamente de $\sqrt{gC}$. Tendo em conta que $g = 10 \text{ ms}^{-2}$ e que o comprimento normal de uma perna é de 0,9 metros, a velocidade máxima para uma caminhada comum será de 3 metros por segundo – uma estimativa bastante boa – e, quanto mais alta for a pessoa, maior será o valor C e mais rápido caminhará, embora devido à raiz quadrada não haja realmente uma grande diferença entre as velocidades de caminhada de pessoas com alturas pouco diferentes.

Outra forma de interpretar este resultado seria observar pessoas (ou outras criaturas com duas pernas) que tentam deslocar-se de um ponto para outro à maior velocidade possível e tentar determinar a que velocidade param de andar e começam a correr. A velocidade crítica $\sqrt{(gC)}$ é o maior ritmo que se pode esperar sem interromper o contacto com o chão (*lifting*, como é conhecido pelos praticantes de marcha atlética). Quando este contacto é interrompido, consegue avançar muito mais depressa, com uma velocidade máxima de $V = \sqrt{(2gnS)} \sim 9\text{ms}^{-1}$, em que $S \sim 0.3\text{ m}$ é a diferença de comprimento entre a perna esticada e a perna fletida quando levanta o pé do chão e $n \sim 10$ é o número de passos que dá para acelerar até à velocidade máxima.

Os praticantes de marcha atlética caminham a uma velocidade muito superior a 3 metros por segundo. O recorde mundial para 1500 metros de marcha, estabelecido pelo americano Tim Lewis em 1988, é de 5 minutos e 13,33 segundos, uma velocidade média de $4,78\text{ms}^{-1}$. Raramente são organizadas competições de marcha para esta distância, portanto será interessante analisar o recorde mundial altamente competitivo dos 20 km de marcha, a mais curta das duas competições. Este foi reduzido para 1 hora, 17 minutos e 16 segundos pelo praticante russo, Vladimir Kanaykin, no dia 29 de setembro de 2007, com uma velocidade média de $4,3\text{ms}^{-1}$ ao longo de mais de 20 km! Estas velocidades conseguem ultrapassar a nossa estimativa de $\sqrt{(gC)}$ por os praticantes de marcha atlética usarem um estilo de caminhada muito mais eficiente do que aquele que usamos quando caminhamos normalmente. Eles não balançam o centro do corpo para cima e para baixo e conseguem virar as ancas de uma forma muito flexível que produz passos muito mais longos e uma frequência de passos muito maior. Este movimento muito eficiente, combinado com uma excelente forma física, permite-lhes manter velocidades impressionantes ao longo de grandes períodos. O detentor do recorde mundial para 50 km tem uma média de mais de $3,8\text{ms}^{-1}$ e consegue percorrer a distância de uma maratona (42,2 km) em 3 horas e 6 minutos.

Todos os números inteiros positivos são amigos pessoais de Ramanuján.

John E. Littlewood

Pense num número entre 1 e 9. Multiplique-o por 9 e some os dois dígitos que compõem o novo número. Subtraia 4 ao resultado e ficará com um número com um só dígito. A seguir converta-o numa letra: se o número for 1 deverá convertê-lo num A, se for 2 num B, 3 num C, 4 num D, 5 num E, 6 num F, etc. Agora pense num animal cujo nome comece pela letra escolhida e imagine esse animal com toda a sua concentração. Mantenha a imagem do animal o mais presente na sua mente possível. Se consultar a nota 14 no fim deste livro, verá que li a sua mente e descobri que animal estava a imaginar.

Este é um truque muito simples e o leitor deverá conseguir descobrir facilmente como adivinhei o animal que escolheu com uma probabilidade de sucesso tão grande. Envolve um pouco de matemática, na medida em que são exploradas algumas propriedades simples dos números, mas contém também um ingrediente psicológico – e zoológico.

Há um outro truque deste tipo que envolve apenas as propriedades dos números. É um truque que usa o número 1089, que talvez já seja um dos seus favoritos. Foi neste ano que houve um terramoto em Inglaterra; também é um quadrado perfeito (33×33); mas a sua propriedade mais espantosa é a que vou explicar a seguir.

Escolha qualquer número de três dígitos, composto por três algarismos diferentes (como 153). Crie um segundo número invertendo a ordem dos três dígitos (no nosso exemplo seria 351). Agora subtraia o menor destes dois números ao maior (ou seja, $351 - 153 = 198$; se o seu número só tiver dois dígitos, como 23, acrescente-lhe um 0 à esquerda, ficando

com 023). Em seguida, some o resultado ao número que obtém se escrever esse mesmo número de trás para a frente ($198 + 891 = 1089$). Qualquer que seja o número escolhido no primeiro passo, acabará a obter o número 1089 depois de seguir esta sequência de operações!¹⁵

54 | O PLANETA DOS ENGANADORES

Pode enganar algumas pessoas durante parte do tempo e outras durante todo o tempo, mas não pode enganar todas as pessoas durante todo o tempo.

Abraham Lincoln

Uma das intuições humanas que foi aperfeiçoada por inúmeras gerações de interação social é a confiança. Baseia-se na nossa capacidade de avaliar a probabilidade de alguém estar a dizer a verdade. Uma das maiores distinções que existem entre diferentes ambientes é o facto de partirmos do princípio de que as pessoas são honestas até prova em contrário ou de acreditarmos que são desonestas até prova em contrário. Encontramos esta distinção na burocracia de diferentes países. Na Grã-Bretanha, todo o sistema judicial e burocrático é baseado na premissa de que as pessoas são honestas, mas pude observar que noutros países se parte do pressuposto contrário e que as regras e regulamentos são criados com base na premissa da desonestidade. Quando fazemos uma participação a uma seguradora, descobrimos rapidamente qual é a postura dela em relação aos seus clientes.

Imagine que na nossa exploração do Universo encontrámos uma civilização no estranho mundo de Janus. A observação da atividade política e comercial deste povo durante um longo período mostrou-nos que, em média, os cidadãos dizem a verdade $\frac{1}{4}$ das vezes e mentem $\frac{3}{4}$ das vezes. Apesar desta média preocupante, decidimos fazer-lhes uma visita e somos recebidos pelo líder do partido maioritário que faz um grande discurso acerca das suas intenções benevolentes. Segue-se um discurso do líder do partido da oposição que diz que a afirmação do líder do partido maioritário era verdadeira. Qual é a probabilidade de a afirmação do líder ser, de facto, verdadeira?

Temos de descobrir qual é a probabilidade de a afirmação do líder do partido maioritário ser verdadeira tendo em conta que o líder da oposição afirmou que o era. Este resultado será igual à probabilidade de a afirmação do líder ser verdadeira, mais a probabilidade de a afirmação do líder da oposição ser verdadeira, dividida pela probabilidade de a afirmação do líder da oposição ser verdadeira. No caso da primeira parte, a probabilidade de ambos estarem a dizer a verdade é apenas de $\frac{1}{4} \times \frac{1}{4} = 1/16$. A probabilidade de o líder da oposição ter dito a verdade é a soma de duas probabilidades: a primeira é a probabilidade de o líder ter dito a verdade e de o líder da oposição também ter dito a verdade, que corresponde a $\frac{1}{4} \times \frac{1}{4} = 1/16$, e a probabilidade de o líder estar a mentir e de o líder da oposição também estar a mentir, que é de $\frac{3}{4} \times \frac{3}{4} = 9/16$. Desta forma, a probabilidade¹⁶ de a afirmação do líder do partido maioritário ser, de facto, verdadeira é apenas de $1/16 / 10/16$, ou seja, $1/10$.

A lotaria é um imposto sobre a tolice. E, graças a Deus, é fácil de cobrar porque a credulidade está sempre na moda.

Henry Fielding

A UK National Lottery (semelhante ao totoloto em Portugal) tem uma estrutura simples. Pagamos 1£ para escolher seis números diferentes da lista de 1, 2, 3, ..., 48, 49. Ganhamos um prêmio se pelo menos três dos números constantes do nosso bilhete corresponderem às seis bolas diferentes selecionadas pela máquina que foi criada para fazer uma escolha aleatória entre as 49 bolas numeradas. Depois de extraídas, as bolas não voltam para dentro da máquina. Quantos mais números acertar, maior o prêmio que recebe. Se acertar nos seis números, partilhará o *jackpot* com todas as outras pessoas que também tenham escolhido os seis números selecionados. Para além destas seis bolas, é tirada uma bola adicional, chamada "número suplementar". Esta afeta apenas os jogadores que já acertaram em cinco dos seis números já tirados. Se também acertarem no número suplementar, recebem um prêmio maior do que aqueles que apenas acertaram nos outros cinco números.

Quais são as suas probabilidades de escolher corretamente os 6 números entre os 49 possíveis, partindo do pressuposto de que a máquina escolhe os números premiados aleatoriamente? A extração de cada bola é um acontecimento independente que não tem qualquer efeito sobre

* Em rigor, há 12 máquinas (cada qual com um nome) e 8 conjuntos de bolas que podem ser usados para a extração pública dos números vencedores na televisão. A máquina e o conjunto de bolas que são utilizados em cada extração são escolhidos aleatoriamente entre estes candidatos. Normalmente este aspeto é ignorado pelas pessoas que elaboram as análises estatísticas dos resultados da extração desde o início da lotaria. Uma vez que a fonte mais provável de um elemento não-aleatório que favoreça um determinado grupo de números em detrimento de outro poderia estar associada a características de uma máquina ou bola em particular, é importante fazer estudos estatísticos separados para cada máquina e conjunto de bolas. Este tipo de tendências seria anulado quando se calculasse a média de todos os conjuntos de bolas e máquinas.

a extração seguinte, exceto na medida em que reduz o número de bolas a partir do qual será feita a seleção. A probabilidade de acertar no primeiro dos 6 números premiados entre os 49 é assim a fração $6/49$. A probabilidade de acertar no seguinte número dos 5 restantes entre as 48 bolas disponíveis passa a ser de $5/48$. A probabilidade de acertar no seguinte número dos 4 restantes entre as 47 bolas disponíveis é de $4/47$, etc., sendo as restantes probabilidades de $3/46$, $2/45$ e $1/44$. Desta forma, a probabilidade de acertar em cada uma destas bolas independentemente e de partilhar o *jackpot* é de:

$$6/49 \times 5/48 \times 4/47 \times 3/46 \times 2/45 \times 1/44 = 720/10\,068\,347\,520$$

Se fizer esta divisão, concluirá que a probabilidade é de 1 em 13 983 816, ou seja, aproximadamente de 1 em 13,9 milhões. Se quiser acertar 5 números mais o suplementar, a probabilidade é 6 vezes menor, e assim a sua probabilidade de partilhar o prémio é de 1 em 13 983 816/6, ou seja, 1 em 2 330 636.

Consideremos o conjunto de todos os resultados possíveis – todos os 13 983 816 – e perguntemos quantos deles resultarão em 5 ou 4 ou 3 ou 2 ou 1 ou zero números escolhidos acertadamente¹⁷. Há apenas 258 pessoas que acertam em 5 números, mas 6 delas ganham o prémio do número suplementar, o que nos deixa com 252; 13 545 acertam em 4 números, 246 820 acertam em 3 números, 1 851 150 acertam em 2 números, 5 775 588 acertam apenas num número e 6 096 454 não acertam em número algum. Assim, para saber qual é a probabilidade de acertar em 5 números, por exemplo, bastará dividir o número de formas pelas quais isso pode acontecer pelo número total de combinações possíveis, ou seja, $252/13\,983\,816$, o que representa uma probabilidade de 1 em 55 491 com um bilhete. As probabilidades de acertar em 4 números são de 1 em 1032; para acertar em 3 bolas tem uma probabilidade de 1 em 57. O número de resultados premiados entre os 13 983 816 possíveis é $1 + 258 + 13\,545 + 246\,820 = 260\,624$, pelo que a probabilidade de ganhar um prémio comprando um bilhete é de 1 em $13\,983\,816/260\,624$, ou seja, cerca de 1 em 54. Se comprar um bilhete por semana e um bilhete adicional no seu aniversário e no Natal, tem uma probabilidade de 50-50 de ganhar algum prémio.

Estes cálculos não são muito animadores. O estatístico John Haigh evidencia que uma pessoa tem mais probabilidades de ter morte súbita após comprar um bilhete da lotaria do que de ganhar o *jackpot*. Embora seja verdade que se não jogar certamente não ganhará, o que acontece se comprar muitos bilhetes?

A única forma de ter a certeza de que ganha a lotaria é comprar *todos* os bilhetes. Já foram feitas várias tentativas de usar esta estratégia em diferentes versões da lotaria em todo o mundo. Se ninguém ganhar o *jackpot* numa determinada extração, normalmente o prémio que não foi atribuído transita para a semana seguinte, criando assim um *jackpot* ainda maior. Nessas situações, poderá ser tentador procurar comprar quase todos os bilhetes – o que é perfeitamente legal! A Virginia State Lottery nos EUA é semelhante à UK Lottery e ao totoloto em Portugal, com a diferença de que os 6 números vencedores são escolhidos entre um conjunto de apenas 44 bolas, o que se traduz em 7 059 052 resultados possíveis. Quando o *jackpot* tinha acumulado o valor de 27 milhões de dólares, o jogador australiano Peter Mandral pôs em funcionamento um esquema bem oleado de compra e impressão que conseguiu comprar 90% dos bilhetes (os 10% que faltam devem-se a uma falha por parte de alguns membros da equipa). Ele ganhou o *jackpot* acumulado e teve um bom lucro em cima dos dez milhões de dólares que gastou na compra dos bilhetes e no pagamento aos seus “colaboradores”.

56 | UM JOGO DE FUTEBOL BIZARRO

Auto-golo: Os *auto-golos*, tal como os *desvios*, têm tendência para ser descritos com piedade por aqueles que são vítimas deles, sendo frequentemente precedidos de adjetivos como *insólito* ou *bizarro* quando as palavras "incompetente" ou "estúpido" seriam as que primeiro nos viriam à mente.

John Leigh e David Wodehouse. *The Football Lexicon*

Qual foi o jogo de futebol mais bizarro de sempre? Nesta competição acho que há apenas um vencedor. Teria de ser o infame jogo entre Granada e Barbados na Taça das Caraíbas em 1994. Este torneio tem uma fase de apuramento por grupos antes das eliminatórias finais. No último jogo da fase de grupos, a equipa de Barbados precisava de ganhar à equipa de Grenada por dois golos, pelo menos, para se qualificar para a etapa seguinte. Se não conseguissem, seria a equipa de Grenada a qualificar-se. Parece muito simples. O que é que poderia correr mal?

Infelizmente, a lei das consequências imprevistas atacou sem piedade. Os organizadores do torneio introduziram uma nova regra para atribuir uma vantagem justa às equipas que ganhassem no prolongamento marcando um "golo de ouro". Uma vez que o golo de ouro punha fim à partida, nunca se podia ganhar por mais do que um golo nestas circunstâncias, o que parecia injusto e prejudicaria uma equipa que precisasse de ganhar por uma vantagem de dois golos. Desta forma, a organização decidiu que o golo de ouro contaria como dois golos. Mas vejamos o que aconteceu.

A equipa de Barbados acabou a primeira parte com uma vantagem de 2-0 e depois fez um jogo menos ofensivo na segunda parte. A apenas sete minutos do fim do jogo, Grenada marcou um golo, ficando a pontuação a 2-1. A equipa de Barbados ainda podia qualificar-se se marcassem um terceiro golo, mas não era fácil tendo em conta que faltavam poucos

minutos para o fim da partida. Decidiram que seria melhor atacarem a própria baliza, criando um empate, pois assim ainda tinham a possibilidade de ganhar com um gol de ouro no prolongamento, que contaria como dois golos, qualificando-se assim às custas da equipa de Grenada. A equipa de Barbados fez isso mesmo e marcou um golo na própria baliza, deixando o jogo a 2-2 a apenas três minutos do fim da partida. Grenada percebeu que se conseguisse marcar outro golo (em qualquer uma das balizas!) passaria para a fase seguinte e, assim, atacaram a própria baliza para marcar o golo que os apuraria para a fase seguinte. Mas a equipa de Barbados defendeu determinadamente o pontapé de Grenada para os impedir de marcar e foram para prolongamento. No prolongamento, a equipa de Barbados apanhou os adversários de surpresa atacando a baliza de Grenada e marcando o golo de ouro da vitória nos primeiros cinco minutos. Se não acredita em mim, veja o vídeo no YouTube*!

* <http://www.youtube.com/watch?v=ThpYsN-4p7w>

O gênio é composto por quatro partes de transpiração e uma parte de perspectiva estratégica.

Armando Iannucci

Um velho arco de pedra pode parecer uma criação muito estranha. Cada pedra parece ter ali sido colocada individualmente, mas a totalidade da estrutura parece não poder ser suportada até que a última pedra tenha sido colocada: não se pode ter um “quase” arco. Então, como são construídos?

Este é um problema interessante por fazer lembrar uma curiosa discussão, muito popular nos Estados Unidos, chamada “*Design Inteligente*”. Em termos gerais, os seus defensores selecionam algumas coisas complicadas que existem no mundo natural e argumentam que devem ter sido “concebidas” com aquela forma em vez de terem evoluído através de um processo gradual a partir de formas mais simples por não haver um passo anterior a partir do qual possam ter-se desenvolvido. Isto é um pouco subjetivo, claro – podemos estar a ser pouco criativos ao imaginar qual poderá ter sido esse passo anterior –, mas, na sua raiz, o problema é um pouco como o nosso arco, que é uma construção complicada que não parece estar um passo mais à frente de uma versão ligeiramente mais simples de um arco com uma pedra em falta.

As conclusões a que chegamos no caso do arco denotam pouca imaginação na medida em que estamos presos a uma linha de raciocínio que determina que todas as estruturas são construídas acrescentando peças. Mas há estruturas que podem ser criadas pela via da subtração. Suponhamos que começámos com um monte de pedras e que as fomos movendo gradualmente e removendo pedras do centro do monte até sobrar apenas um arco. Se encararmos a estrutura por este ponto de vista, percebemos

qual é o aspeto do “quase” arco. Parte da abertura central está preenchida. Os verdadeiros arcos subaquáticos são formados pela erosão gradual do buraco ao centro até que resta apenas o arco exterior. Da mesma forma, nem toda a complexidade do mundo natural é formada através da soma.



O Nobre Caminho Óctuplo: Visão correta, intenção correta, linguagem correta, ação correta, modo de vida correto, esforço correto, atenção correta e concentração correta.

O Nobre Caminho Óctuplo

Contamos em “dezenas”. Dez unidades fazem dez, dez dezenas fazem cem, dez centenas fazem mil, etc. É por isso que o nosso sistema de contagem se chama “decimal”. O seu alcance é ilimitado desde que haja nomes suficientes para designar os resultados. Referimo-nos a milhões e milhares de milhões para designar alguns dos números grandes, mas não temos nome para todos os números que possamos precisar de escrever. Em vez disso, temos uma noção muito prática que consiste em escrever 10^n para denotar o número que é 1 seguido de n zeros, da mesma forma que mil é igual a 10^3 .

Não é difícil encontrar a origem de todas estas dezenas no sistema de contagem. Está na ponta dos nossos dedos. A maioria das antigas culturas humanas usava os dedos de alguma forma na contagem. Consequentemente, encontramos sistemas de contagem baseados em grupos de 5 (os dedos de uma mão), 10 (os dedos das duas mãos), de 20 (os dedos das mãos e dos pés) ou misturas de todos ou alguns destes três sistemas. O nosso sistema de contagem denuncia uma história complicada em que diferentes sistemas se fundiram para dar origem a novos sistemas pela presença das palavras antigas que refletem a base anterior. Assim, temos a palavra “dúzia” para designar o número 12, e em inglês existe ainda a palavra *score* (derivada da palavra saxónica *sceran*, que significava cortar ou marcar com entalhes), que também designa o número 20.

Apesar da ubiquidade do sistema de contagem baseado em dezenas nas culturas mais antigas, existe um caso invulgar em que uma sociedade

indígena da América Central usava um sistema de contagem baseado no número 8. Consegue imaginar por que motivo escolherem este número? Eu costumava perguntar aos matemáticos se conseguiam encontrar uma boa explicação e normalmente respondiam-me que 8 era um bom número a usar por ter muitos fatores, por ser exatamente divisível por 2 e 4, permitindo dividir porções em quartos sem ter de criar um novo tipo de quantidade a que chamamos fração. A única altura em que obtive a resposta certa foi quando perguntei a um grande grupo de crianças de 8 anos e uma rapariga me deu imediatamente a resposta certa: estavam a contar os espaços entre os dedos. Se segurarmos objetos nos espaços entre os dedos, como pedaços de fio ou de tecido, esta é uma forma de contar natural. Os autores do sistema de contagem baseado no número 8 também estavam a contar pelos dedos.

59 | UMA MAIORIA REPRESENTATIVA

A democracia costumava ser uma coisa boa, mas agora caiu nas mãos erradas.

Jesse Helms

Os políticos têm o hábito de presumir que têm uma representatividade muito maior do que realmente têm. Se tivermos uma lista de políticas a oferecer a um eleitorado, o facto de ter conseguido a maior parte dos votos não significa que a maioria dos eleitores favoreça cada uma das suas políticas em detrimento da alternativa da oposição. E se tiver ganho as eleições por uma margem estreita, que tipo de representatividade tem realmente?

Por uma questão de simplicidade, suponhamos que há apenas dois candidatos (ou partidos) a participar nas eleições. Suponhamos que o vencedor teve V votos e que o perdedor teve P votos, sendo o número total de votos válidos igual a $V + P$. Em qualquer número de acontecimentos com esta dimensão, o "erro" estatístico aleatório que se espera que ocorra é determinado pela raiz quadrada de $V + P$, pelo que se $V + P = 100$, haverá uma incerteza estatística de 10 para qualquer um dos lados. Para o vencedor das eleições poder ter a certeza de que não ganhou devido a uma variação aleatória significativa no processo de votação – como a contagem e o tratamento dos votos – temos de ganhar por uma margem superior à variação aleatória:

$$V - P > \sqrt{(V+P)}$$

Se deram entrada nas urnas 100 votos, a margem de vitória tem de exceder os 10 votos para ser convincente. A título de exemplo, nas eleições

presidenciais nos EUA em 2000*, Bush recebeu 271 votos do Colégio Eleitoral e Gore recebeu 266. A diferença foi apenas de 5, muito inferior à raiz quadrada de $271 + 266$, que é aproximadamente 23.

Uma história ainda mais divertida conta que Enrico Fermi, o grande físico italiano que teve um papel fundamental na criação da primeira bomba atômica – e que era um jogador de ténis muito competitivo – certa vez, ao ser derrotado no ténis por 6 jogos contra 4, respondeu que a derrota não era estatisticamente significativa por a margem ser inferior à raiz quadrada do número de jogos!

Supondo que o leitor ganhou as eleições e teve uma margem de vitória suficientemente grande para dissipar as dúvidas de que se tenha devido puramente a erros aleatórios, que maioria necessita de ter a seu favor para poder afirmar que tem uma representatividade suficiente para garantir que o eleitorado apoia as suas políticas? Uma sugestão interessante seria exigir que a fração dos votos que o vencedor obteve, $V/(V+P)$, seja superior à proporção dos votos do perdedor em relação aos do vencedor, P/V . Esta condição para uma representatividade “dourada” traduz-se na seguinte fórmula:

$$V/(V+P) > P/V$$

Esta condição exige que $V/P > (1+\sqrt{5})/2 = 1,61$, que é a famosa “proporção dourada”. Isto significa que o leitor necessitaria de obter uma fração $V/(V+P)$ maior do que $8/13$ ou 61,5% dos votos totais nos dois partidos. Nas últimas eleições gerais no Reino Unido, o Partido Trabalhista ganhou 412 assentos parlamentares e o Partido Conservador 166, pelo que os Trabalhistas conseguiram 71,2% destes assentos, o suficiente para uma representatividade “dourada”. Pelo contrário, nas eleições de 2004 nos EUA, Bush teve 286 votos do colégio eleitoral, Kerry teve 251, e assim Bush recebeu apenas 53,3% do total dos votos, menos do que é exigido para uma representatividade “dourada”.

* Alguns aspetos desta eleição continuam a ser considerados muito duvidosos do ponto de vista estatístico. Na crucial votação da Flórida, o resultado da recontagem foi muito misterioso. Na recontagem dos votos foram encontrados 2200 votos para Gore e 700 para Bush. Uma vez que se calculava que houvesse uma probabilidade igual de os votos ambíguos poderem pertencer a um ou outro candidato, a grande assimetria na atribuição dos votos na recontagem sugere que se passou algo de natureza não-aleatória na contagem inicial ou na recontagem.

60 | O CAMPEONATO DE DUAS CABEÇAS

Os primeiros serão os últimos; e os últimos serão os primeiros.

Evangelho de S. Mateus

Em 1981, a Associação de Futebol de Inglaterra fez uma mudança radical no sistema de pontuação dos seus campeonatos, numa tentativa de recompensar um tipo de jogo mais ofensivo. Propuseram que fossem atribuídos 3 pontos às vitórias em vez dos 2 pontos que eram tradicionalmente atribuídos à equipa vencedora. Os empates continuavam a receber 1 ponto. Outros países em breve seguiram-lhes o exemplo e este é agora o sistema universal de pontuação nas competições das ligas de futebol em todo o mundo. É interessante observar o efeito que esta alteração teve no nível de sucesso que uma equipa desinteressante e não vencedora pode alcançar. No tempo em que eram atribuídos 2 pontos às vitórias, era possível ganhar facilmente o campeonato com 60 pontos em 42 jogos, pelo que uma equipa que alcançasse os 42 pontos por ter empatado em todos os jogos podia acabar o campeonato nos primeiros lugares – e, de facto, o Chelsea ganhou o campeonato da Primeira Divisão em 1955 com a pontuação mais baixa de sempre: 52. Hoje em dia, com 3 pontos a serem atribuídos a cada vitória, a equipa campeã precisa de ter mais de 90 pontos em 38 jogos e uma equipa que tivesse empatado em todos os jogos descobriria que os 42 pontos alcançados com o seu desempenho a deixariam no terceiro ou quarto lugar a contar do fim, a lutar para não descer de divisão.

Com estes argumentos em mente, imaginemos uma liga cujas autoridades decidissem mudar o sistema de pontuação imediatamente após o último jogo da temporada. Ao longo de toda a época, as equipas tinham estado a jogar para obter 2 pontos por uma vitória e 1 ponto

por um empate. O campeonato é jogado por 13 equipes e jogam uma vez contra cada uma das outras, pelo que cada equipe joga 12 partidas. Os All Stars ganham 5 jogos e perdem 7. Incrivelmente, todos os outros jogos do campeonato resultam num empate. Desta forma, os All Stars obtêm um total de 10 pontos. Todas as outras equipes ganham 11 pontos com os seus 11 empates e 7 delas ganham 2 pontos adicionais ao ganharem aos All Stars, ao passo que 5 delas perdem para os All Stars e não ganham pontos nesses jogos. Assim, 7 das outras equipes acabam o campeonato com 13 pontos e 5 delas acabam com 12 pontos. Todas as equipes têm mais pontos que os All Stars, que se veem assim atirados para o fundo da tabela das classificações do campeonato.

Quando os All Stars, desanimados, regressam ao balneário depois do jogo final e percebem que estão no último lugar da tabela de classificações, enfrentando a descida de divisão e a provável ruína financeira, surge a notícia de que as autoridades da liga decidiram introduzir um novo sistema de pontuação que será aplicado retroativamente a todos os jogos daquela temporada. Para recompensar o jogo ofensivo, decidem atribuir 3 pontos a cada vitória e 1 ponto a cada empate. Os All Stars apressam-se a refazer os cálculos. Agora obtiveram 15 pontos com as suas 5 vitórias. As outras equipes continuam a obter 11 pontos com os seus 11 empates. Mas agora, os 7 que venceram os All Stars só recebem 3 pontos cada, ao passo que os 5 que perderam contra eles não recebem nenhum ponto. De qualquer das formas, as outras equipes terminam o campeonato com 11 ou 14 pontos e os All Stars são agora os campeões!

61 | CRIAR ALGO A PARTIR DO NADA

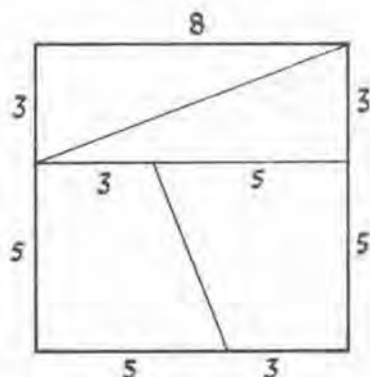
Os erros são bons. Quanto mais erros, melhor. As pessoas que cometem erros são promovidas. São de confiança. Porquê? Não são perigosas. Não podem ser assim tão sérias. As pessoas que não cometem erros acabam a cair de penhascos, o que é uma coisa má porque todas as pessoas em queda livre são consideradas um risco. Podem cair em cima de nós.

James Church, *A Corpse in the Korymbos*

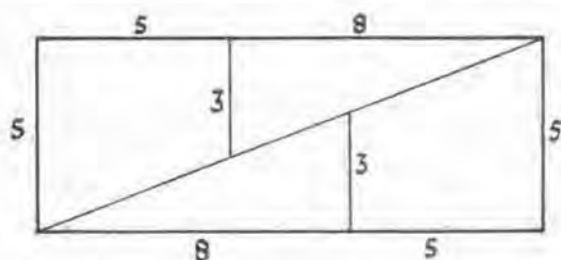
Se o leitor for uma daquelas pessoas que dão palestras ou fazem “apresentações” com programas de computador como o PowerPoint, provavelmente também já descobriu um dos pontos fracos deste sistema – especialmente se for professor. Quando acaba a apresentação, normalmente chega a altura reservada para perguntas do público acerca do que acabou de dizer. Uma característica dessas perguntas é serem frequentemente mais fáceis de responder por escrito ou por meio de um desenho. Se estiver a usar um quadro negro ou um projetor com folhas de acetato e tiver uma caneta, não terá dificuldade em desenhar ou escrever a explicação necessária. Mas se estiver armado apenas com o computador, dá consigo num beco sem saída. Não pode “desenhar” facilmente algo em cima da apresentação, a não ser que tenha um “tablet”. Isto serve apenas para demonstrar o quanto dependemos das imagens para explicar o que dizemos. São mais diretas do que as palavras. São analógicas em vez de digitais.

Alguns matemáticos não confiam em desenhos. Gostam de ver provas que não façam referência a nenhuma representação gráfica que possa ter sido desenhada e que influencie a nossa ideia da verdade. No entanto, a maioria dos matemáticos é da opinião contrária. Gostam de imagens e consideram-nas um guia essencial para perceber o que pode ser verdade

e como essa verdade pode ser demonstrada. Uma vez que esta é a opinião da maioria, vejamos algo que deixará a minoria muito contente. Suponhamos que temos 8×8 metros quadrados de um revestimento caro para o chão que é composto por quatro peças – dois triângulos e dois quadriláteros – como se pode ver na representação abaixo.



É fácil constatar que a área total do quadrado é de $8 \times 8 = 64$ metros quadrados. Mas peguemos nos nossos quatro pedaços de revestimento com as dimensões indicadas e disponhamo-los de forma diferente. Desta vez criamos um retângulo, como se vê na representação seguinte:



Aconteceu algo estranho. Qual é a área do novo revestimento retangular? É $13 \times 5 = 65$ metros quadrados¹⁸. Criámos um metro quadrado de revestimento a partir do nada! O que é que aconteceu? Experimente recortar pedaços com estas dimensões e faça a experiência.

62 | COMO MANIPULAR UMA ELEIÇÃO

Aceito ser consultor do seu grupo nas próximas eleições. Digam-me quem querem que ganhe. Depois de falar com os membros, desenvolverei um "procedimento democrático" que garanta que o vosso candidato ganha as eleições.

Donald Saari¹⁰

Como já vimos no Capítulo 59, as eleições podem ser uma coisa complicada. Há muitas formas de contar os votos e se isto não for feito de forma sensata podemos descobrir que o candidato A ganha ao candidato B, que ganha ao candidato C, que perde para o candidato A! Este resultado não é desejável. Às vezes damos connosco a votar num conjunto de candidatos, em que o mais fraco é eliminado em cada fase, o que resulta na transferência dos votos desse candidato para os restantes candidatos que passam para a etapa seguinte da eleição.

Mesmo que não passe os seus dias a eleger candidatos, ficará surpreendido com a quantidade de vezes que é envolvido em votações. Que filme devemos ir ver? Que canal de TV devemos escolher ver? Onde devemos ir passar as férias? Qual é a melhor marca de frigorífico a comprar? Se discutir com outras pessoas questões que admitem várias respostas diferentes, estará a participar numa "votação" – que expressa a sua preferência – e o candidato vitorioso é a escolha obtida. Nem todas as decisões são alcançadas por meio de uma votação organizada em que todos participam. Normalmente, o processo é muito menos rigoroso. Alguém sugere um filme. Depois alguém sugere outro filme porque é mais recente. Depois alguém diz que o filme novo é demasiado violento e que deveríamos escolher um terceiro. Em seguida descobre-se que uma das pessoas já viu o terceiro filme sugerido, portanto volta-se à primeira opção. Depois alguém percebe que esse filme não é adequado para crianças, portanto sugerem outro. Por esta

altura, as pessoas já estão fartas e concordam com esta última proposta. O que está acontecer neste cenário é muito interessante. Estamos a considerar uma possibilidade de cada vez, em comparação com uma alternativa, e o processo é repetido como rondas de um torneio. Não escolhemos analisar todas as características de todos os filmes possíveis em conjunto e fazer uma votação. Desta forma, o resultado das deliberações depende fortemente da ordem pela qual um filme é considerado em alternativa a outro filme. Bastará mudar a ordem pela qual os filmes são considerados e os critérios usados para os comparar e o resultado pode ser muito diferente.

Com as eleições acontece exatamente o mesmo. Suponhamos que temos 24 pessoas que têm de escolher um líder entre 8 candidatos possíveis (A, B, C, D, E, F, G e H). Os “eleitores” dividem-se em três grupos que optam pelas seguintes ordens de preferência dos candidatos:

1.º grupo: A B C D E F G H

2.º grupo: B C D E F G H A

3.º grupo: C D E F G H A B

À primeira vista parece que o candidato C é o preferido por todos os grupos, ficando em primeiro, segundo e terceiro lugar nas três listas de preferências. No entanto, a mãe do candidato H quer muito que ele seja eleito e pergunta-nos se não podemos arranjar maneira de garantir que ele vence a eleição. Parece impossível, tendo em conta que ele está em último, penúltimo e antepenúltimo lugar nas listas de preferência. Não parece haver a menor possibilidade de ele vir a ser líder. Explicamos à mãe de H que o procedimento tem de seguir as regras e que não pode haver desonestidade. Assim, o desafio é encontrar um sistema de votação que permita que H vença.

Para o conseguir, basta organizar uma eleição estilo torneio, e escolher o vencedor entre cada grupo de 2 adversários usando as preferências dos três grupos indicados acima. Primeiro, colocamos G contra F e vemos que F ganha 3-0. A seguir, F defronta-se com E e perde 3-0. Depois E concorre com D e perde 3-0. D compete com C e perde 3-0. C compete com B e perde 2-1. Depois B compete com A e perde 2-1, o que deixa A contra H na última ronda. H vence A por 2-1. Desta forma, H é a escolha vencedora deste “torneio” que foi criado para escolher o novo líder.

Qual foi o truque? Basta certificarmo-nos de que os candidatos mais fortes se eliminam uns aos outros nas primeiras etapas e só incluir o nosso candidato "protegido" no último momento, assegurando que só é comparado com candidatos que pode vencer. Desta forma, se escolhermos cuidadosamente a ordem das rondas, um tenista britânico pode ganhar o torneio de Wimbledon.

63 | O BALANÇO DO PÊNDULO

A Inglaterra balança como um pêndulo.

Roger Miller

Conta-se que, no século XVI, o grande cientista italiano Galileu Galilei costumava divertir-se a ver balançar o grande candelabro de bronze que pendia do teto da catedral de Pisa. Podia ter sido posto em movimento para espalhar o aroma do incenso ou podia ter sido perturbado pela necessidade de ser baixado para substituir as velas. O que via fascinava-o. A corda que suspendia o candelabro do teto era muita longa, pelo que ele balançava para trás e para a frente como um pêndulo, muito lentamente: suficientemente devagar para permitir calcular o tempo que demoraria a deslocar-se a um dos extremos da sua oscilação e a regressar ao ponto de partida. Galileu observou o que acontecia em muitas ocasiões. De cada vez o candelabro balançava de forma diferente; por vezes fazia oscilações muito pequenas; outras vezes eram oscilações um pouco maiores. Ainda assim, observou algo muito importante. O período de tempo necessário para o candelabro se deslocar até ambos os extremos do movimento pendular era o mesmo, independentemente da distância percorrida. Se fosse empurrado com mais força, ia mais longe do que se fosse empurrado com pouca força, mas quando ia mais longe, percorria essa distância mais depressa e demorava o mesmo tempo a regressar ao ponto de partida que demoraria se tivesse sofrido apenas um ligeiro empurrão.

As consequências desta descoberta* são muito importantes. Se tiver um relógio de pêndulo, tem de lhe dar corda aproximadamente uma vez

* Galileu pensou que isto era verdade em relação a todas as oscilações do candelabro, independentemente da distância percorrida, o que não corresponde exatamente à realidade. É verdade, com um elevado grau de exatidão, para oscilações de "pequena" amplitude. Este tipo de oscilação é conhecido pelos cientistas como "movimento harmónico simples". Descreve o comportamento de quase todos os sistemas estáveis na Natureza depois de o seu estado de equilíbrio ter sido ligeiramente perturbado.

por semana. A descoberta de Galileu significa que, se o pêndulo estiver parado, não importa a forma como o seu impulso o coloca novamente em movimento. Desde que não dê um impulso demasiado forte, o pêndulo demorará o mesmo tempo a deslocar-se para a frente e para trás e o tique-taque daí resultante terá a mesma duração. Se assim não fosse, os relógios de pêndulo seriam objetos muito aborrecidos. Teríamos de definir a amplitude do balanço do pêndulo com muita precisão para o relógio marcar a mesma hora que marcava antes de parar. De facto, esta observação muito precisa de Galileu esteve na origem da criação do relógio de pêndulo. A primeira versão operacional deste tipo de relógio foi construída pelo físico holandês Christiaan Huygens na década de 1650.

Finalmente, há um teste interessante baseado na oscilação de um pêndulo para determinar se a lógica de um físico é mais forte do que o seu instinto de sobrevivência. Imagine que tem um pêndulo muito grande e pesado, como o candelabro da catedral que Galileu observou. Posicione-se de um dos lados do pêndulo e puxe o peso dele na sua direção até estar quase a tocar no seu nariz de físico. Agora solte o peso. Não empurre, limite-se a soltá-lo. Ele vai afastar-se na direção contrária e depois voltar a recuar até tocar na ponta do seu nariz. Quando isso acontecer vai estremecer? Deve estremecer? A resposta é, respectivamente "sim" e "não".

* O pêndulo não pode recuar a uma altura superior àquela a que iniciou o seu movimento (a não ser que alguém o empurre para lhe dar mais energia). Na prática, o pêndulo perde sempre um pouco de energia por causa da resistência do ar e da fricção no seu ponto de apoio, pelo que nunca recuará até exatamente a mesma altura de onde partiu. O físico estará sempre em segurança, mas isso não o impedirá de estremecer.

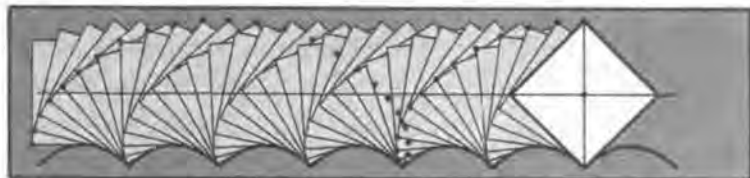
64 | UMA BICICLETA COM RODAS QUADRADAS

A bicicleta é uma companhia tão boa como a maioria dos maridos e, quando fica velha e gasta, a mulher pode livrar-se dela e arranjar uma nova sem chocar toda a comunidade.

Ann Strong

Se a sua bicicleta for como a minha, tem rodas circulares. Pode ter só uma, embora provavelmente tenha duas, mas são certamente circulares. No entanto, o leitor ficará surpreso se eu lhe disser que não tem de ser assim. É perfeitamente possível andar numa bicicleta com rodas quadradas, desde que o faça numa superfície adequada.

Para um ciclista, a característica mais importante de uma roda em movimento é que não dê solavancos enquanto a bicicleta avança. Isso acontece quando uma bicicleta com rodas redondas se desloca sobre uma superfície lisa: o centro de gravidade do ciclista avança em linha reta quando a bicicleta se desloca também em linha reta sem deslizar. Se usar rodas quadradas numa superfície lisa, fará uma viagem muito desconfortável. Mas haverá um tipo de piso diferente que resulte numa viagem confortável usando rodas quadradas? Para responder a esta pergunta, precisamos apenas de descobrir se existe uma forma que produza um movimento em linha reta usando rodas quadradas.



A resposta é muito surpreendente. A forma da superfície da estrada que produz um andar estável com rodas quadradas é criada quando se penduram dois extremos de uma corrente à mesma altura acima do chão. Esta forma chama-se catenária e já falámos dela no Capítulo 11. Se virarmos uma catenária de cabeça para baixo, obtemos uma forma que é usada em muitos dos grandes arcos que existem no mundo. Mas se pegarmos num arco catenário e repetirmos essa forma em sucessão ao longo de uma linha, obtemos uma sequência de ondulações com a mesma altura. É esta a forma da superfície ideal para andar numa bicicleta com rodas quadradas. Basta que os vértices dos quadrados encaixem nos sucessivos "vales" da superfície da estrada. O traço principal da catenária é que quando são colocadas umas a seguir às outras, o ângulo que se forma entre os arcos adjacentes, quando chegam ao ponto mais baixo do vale, é um ângulo reto, com 90 graus, que é também o ângulo de cada um dos vértices do quadrado. Desta forma, uma roda em ângulo reto mantém-se em rotação permanente*.

* O quadrado não é a única forma de roda que permite um passeio sem solavancos. Qualquer roda com uma forma poligonal funcionará numa estrada diferente com forma de catenária. À medida que o número de lados do polígono aumenta e se torna muito grande, começa a parecer-se cada vez mais com um círculo e a linha de catenárias vai ficando progressivamente mais lisa até se parecer com uma estrada perfeitamente plana. O caso de um polígono com três lados, uma roda triangular, é problemático porque chega ao lado da catenária seguinte antes de conseguir encaixar o canto na depressão que separa as catenárias. É preciso alinhar as catenárias da estrada uma a uma para evitar estas colisões. Para uma roda poligonal em movimento com N lados iguais (em que $N = 4$ no caso da roda quadrada), a equação para a estrada com a forma de catenária que permite um passeio regular e em linha recta ao ciclista é $y = -B \cosh(x/B)$, em que $B = C \cot(\pi/N)$ sendo C uma constante.

65 | DE QUANTOS GUARDAS PRECISA UMA GALERIA DE ARTE?

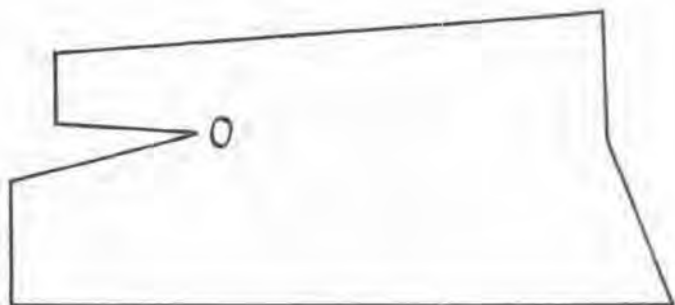
Quem guardará os guardas?

Juvenal

Imagine que é o responsável pela segurança numa grande galeria de arte. Tem muitos quadros valiosos em exposição. Esses quadros estão pendurados num ponto muito baixo, de modo a poderem ser vistos com facilidade, o que os torna vulneráveis ao roubo ou ao vandalismo. A galeria é composta por um conjunto de diferentes salas com formas e tamanhos distintos. Como fará para se certificar de que cada um dos quadros pode ser mantido constantemente sob vigilância? A solução é simples se o seu dinheiro for ilimitado: basta colocar um guarda junto de cada quadro. Mas as galerias de arte raramente nadam em dinheiro e os mecenas ricos não têm o hábito de destinar as suas doações à provisão de guardas e respetivas cadeiras. Portanto, na prática, temos um problema matemático: qual é o menor número de guardas que precisamos de contratar e onde devem ser colocados para que todas as salas da galeria sejam vigiadas?

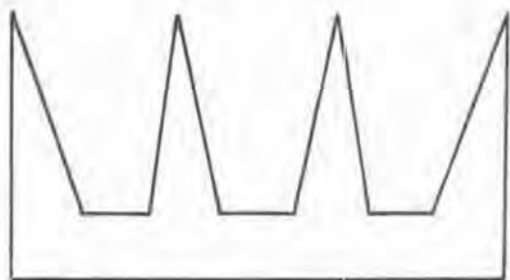
Precisamos de calcular o número de guardas necessários para guardar N paredes. Partimos do pressuposto de que as paredes são direitas e que um guarda no canto onde duas paredes se unem conseguirá ver a totalidade da superfície dessas duas paredes; e supomos também que não há nada a obstruir a visibilidade do guarda. Uma galeria quadrada pode, obviamente, ser vigiada por um único guarda. De facto, se a galeria tiver a forma de qualquer polígono com todos os lados voltados para fora (um polígono "convexo"), um guarda será sempre suficiente.

As coisas tornam-se mais interessantes quando as paredes não estão todas voltadas para fora. Aqui está o exemplo de uma galeria deste tipo, com 8 paredes, que também pode ser vigiada por um único guarda situado no canto O.



Assim, a segurança desta galeria é especialmente económica.

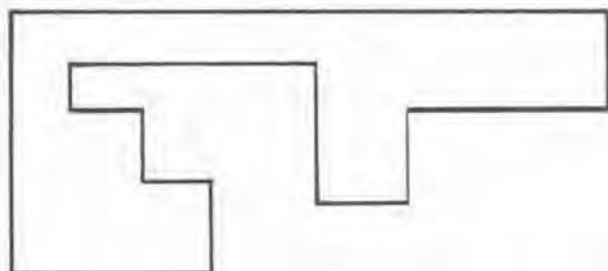
O exemplo seguinte mostra uma galeria mais bizarra, com 12 paredes, que já é menos eficiente. São necessários quatro guardas para vigiar todas as paredes da galeria.



De um modo geral, para resolver este problema basta descobrir como dividir a galeria em triângulos que não se sobreponham – o que é sempre possível. Se o polígono tiver S vértices, originará $S-2$ triângulos. Uma vez que o triângulo é uma dessas formas convexas (com três lados) que necessitam de apenas um guarda, sabemos que se a superfície total da galeria puder ser totalmente coberta por, digamos, T triângulos não sobrepostos, esta poderá sempre ser vigiada por T guardas. Pode, obviamente, ser vigiada por menos guardas. Por exemplo, podemos dividir um quadrado em dois triângulos, unindo as diagonais opostas, mas não precisamos de

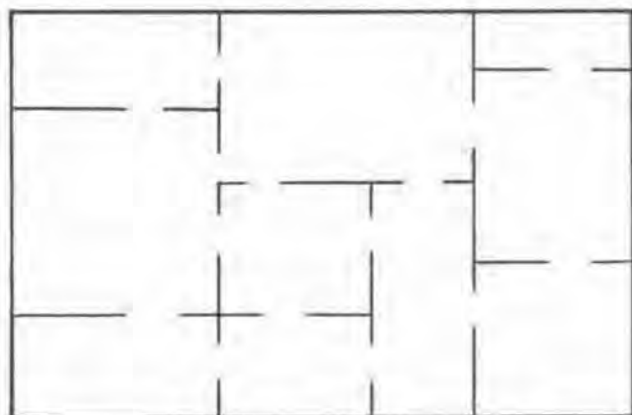
dois guardas para vigiarem as paredes de um quadrado – um guarda é suficiente. De um modo geral, o número de guardas necessários para vigiar uma galeria com P paredes é o número que compõe a parte inteira de $P/3$. Para a nossa galeria com 12 paredes e com a forma de pente, esse número máximo corresponde a $12/3 = 4$, ao passo que para uma galeria com 8 lados, esse número é 2. Infelizmente, determinar se precisamos de usar o número máximo de guardas já não é tão fácil e é aquilo a que se chama um problema “difícil” em informática (ver o Capítulo 27), cujo tempo de cálculo pode ser duplicado de cada vez que se acrescenta mais uma parede ao problema.

A maioria das galerias que o leitor visitará não tem plantas irregulares e denteadas como as destes exemplos. Terão paredes que formam ângulos retos, como esta:



Se existirem muitos cantos numa galeria poligonal formada por ângulos retos como a do exemplo acima, o número de guardas que será necessário colocar nos cantos e que será sempre suficiente para vigiar toda a galeria é a parte inteira de $\frac{1}{4} \times$ (número de cantos). Para a galeria com 14 cantos aqui apresentada, esse número é 3. Isto significa que é muito mais económico ter uma galeria com uma planta deste tipo, no que diz respeito ao dinheiro gasto em ordenados, especialmente se a galeria for grande. Se tiver 150 paredes, a planta não composta por ângulos retos pode necessitar de 50 guardas, ao passo que uma planta composta por ângulos retos necessitará no máximo de 37.

Outro tipo tradicional de galeria composta por ângulos retos é a que é dividida em muitas salas, como neste exemplo com 10 salas:



Nestes casos, pode sempre dividir-se a galeria num conjunto de retângulos que não se sobrepõem. É um tipo de disposição muito útil porque se colocarmos um guarda na porta que liga duas salas, ele pode guardar ambas ao mesmo tempo. Mas nenhum guarda pode vigiar 3 ou mais salas em simultâneo. Assim, o número de guardas suficiente para vigiar completamente a galeria é a parte inteira de $\frac{1}{2} \times (\text{número de salas})$ para a galeria aqui apresentada. Este é um uso muito económico dos recursos humanos.

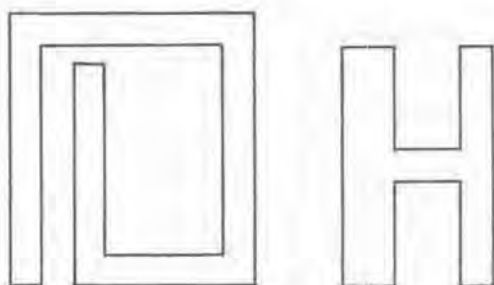
Até aqui falámos de pessoas a vigiar paredes, mas o mesmo exemplo aplica-se à colocação de câmaras de videovigilância ou de lâmpadas destinadas a iluminar a galeria e todas as suas salas. Da próxima vez que o leitor planear roubar a *Mona Lisa*, já tem metade do trabalho feito.

Todos os meus contactos com criminosos me mostraram que o que eles estavam a fazer era apenas uma versão um pouco mais extrema do que todas as outras pessoas fazem.

David Carter

As galerias de arte não são o único tipo de edifícios que precisa de ser vigiado. As prisões e os castelos também necessitam de guardas. No entanto, estes são a versão oposta das galerias de arte – o que é necessário vigiar nestes casos são as paredes exteriores. Quantos guardas têm de ser posicionados nos cantos de uma fortaleza poligonal para vigiarem a total extensão das suas paredes exteriores? A resposta é simples: o número inteiro mais pequeno que seja pelo menos igual a $\frac{1}{2} \times$ (número de cantos). Assim, com 11 cantos, precisamos de 6 guardas para vigiarem as paredes exteriores. O que é ainda melhor é que sabemos que este é *exatamente* o número necessário. Menos de 6 não serão suficientes e mais não serão necessários. No problema que se referia ao interior da galeria, só conhecíamos o maior número de que poderíamos precisar.

Podemos voltar a considerar o caso da prisão com as paredes em ângulos retos seguindo o raciocínio que usámos para as galerias de arte. Podemos ter uma parede exterior na prisão com estas duas formas em ângulos retos.



Nestes dois casos em que as paredes são dispostas de modo a formar ângulos retos, precisamos de 1 mais o menor número inteiro que seja igual a $\frac{1}{4} \times$ (número de cantos) para vigiar a totalidade da parede exterior. Não é possível fazê-lo com menos e não são necessários mais. Nos dois exemplos aqui apresentados há 12 cantos, pelo que precisamos de $1 + 3 = 4$ guardas.

67 | UM TRUQUE DE SNOOKER

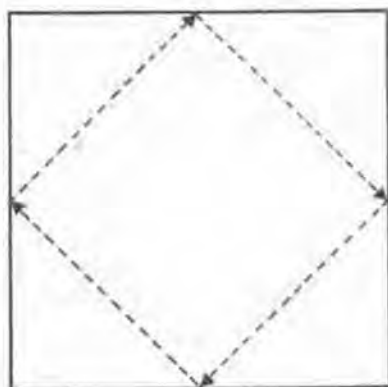
Steve vai atacar a bola cor-de-rosa – e para aqueles que estão a ver a preto e branco, a cor-de-rosa é a que está ao lado da verde.

Ted Lowe

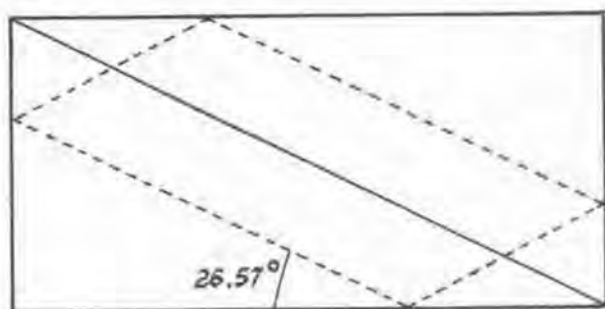
Algumas pessoas sentem-se satisfeitas por os seus filhos passarem a maior parte do tempo a jogar computador, porque isso os ajuda a desenvolver os seus conhecimentos de matemática e de informática. Sempre me perguntei se achavam que o tempo passado a jogar *snooker* ou bilhar aumentava os seus conhecimentos de mecânica newtoniana. Ainda assim, com alguns conhecimentos básicos de geometria conseguirá impressionar qualquer amador.

Suponhamos que quer atingir uma única bola, de modo a fazê-la dar a volta à mesa, bater nas três tabelas e regressar ao ponto de partida. Começemos pelo caso mais simples – uma mesa quadrada. Tudo é convenientemente simétrico e parece óbvio que deve colocar a bola a meio de um dos lados da mesa e depois fazê-la deslizar para o lado, formando um ângulo de 45 graus em relação ao lado da mesa. Ela irá bater no meio do lado adjacente com o mesmo ângulo e seguirá um percurso que forma um quadrado perfeito, apresentado a tracejado no diagrama seguinte.

Claro que não precisa de começar com a bola encostada a um dos lados da mesa; se a atingir em qualquer ponto do quadrado apresentado a tracejado no diagrama e a lançar ao longo de um dos lados do quadrado tracejado, a bola acabará por regressar ao ponto de partida (desde que seja atingida com força suficiente). Se quiser que ela pare exatamente no ponto em que foi atingida, precisa de muita habilidade – ou, pelo menos, de muita prática.



Infelizmente não é comum encontrarmos mesas de *snooker* quadradas. As mesas modernas são do tamanho de dois quadrados colocados lado a lado e uma mesa normal tem habitualmente 3,5 m x 1,75 m. A principal característica destas mesas é o comprimento ser igual ao dobro da largura. Conhecendo este facto simples, podemos desenhar a forma como a tacada teria de ser executada numa mesa retangular com estas dimensões. Incluí no desenho a diagonal para referência. A tacada tem de ser paralela à diagonal e bater nos lados da mesa em pontos que dividam cada um deles numa proporção de 2 para 1, a mesma proporção do comprimento em relação à largura da mesa. (No caso da mesa quadrada, esta proporção era igual a 1 e a tacada tinha de bater no meio de cada um dos lados da mesa.) Isto significa que o ângulo que a bola forma no seu percurso em relação aos lados mais longos da mesa tem uma tangente igual a $\frac{1}{2}$ ou 26,57 graus e o ângulo que forma em relação aos lados mais curtos é 90 menos este ângulo, ou 63,43 graus, uma vez que os três ângulos interiores de um triângulo retângulo têm de somar 180 graus. O paralelogramo assinalado a tracejado marca o único caminho na mesa retangular que fará a bola regressar ao ponto de partida.



Se jogar numa mesa com dimensões diferentes das regulamentares, terá de refazer estes cálculos. De um modo geral, o ângulo mágico a que precisa de atingir a bola em relação ao lado mais longo da mesa é o ângulo cuja tangente é igual à proporção entre a largura e o comprimento da mesa (1:2 no caso da nossa mesa retangular e 1:1 no caso da mesa quadrada), e o ponto da tabela onde precisa de dar a tacada tem de dividir o comprimento da mesa na mesma proporção.

A solidariedade feminina é poderosa.

Robin Morgan

Uma das coisas mais estranhas na China são as consequências crescentes da política do "filho único". Com a exceção dos nascimentos de gêmeos (tipicamente 1% do total), todos os jovens das áreas urbanas são filhos únicos*. Nesta situação, a probabilidade de cada um dos filhos receber excesso de atenção por parte dos pais deu origem à expressão "Síndrome do Pequeno Imperador". No futuro, não haverá irmãos e irmãs nem tios e tias na realidade de praticamente ninguém. Conceitos como "fraternidade" acabarão por perder completamente o sentido.

A princípio parece haver uma estranha assimetria entre os irmãos e irmãs. Se houver 2 filhos, 1 rapaz e 1 rapariga, o rapaz tem 1 irmã, mas a rapariga não tem nenhuma. Se forem 4 filhos, 3 raparigas e 1 rapaz, o rapaz tem 3 irmãs e as raparigas têm entre si $3 \times 2 = 6$ irmãs. Cada uma das raparigas só pode contar as outras duas como irmãs, mas o rapaz conta-as às três. Segundo este ponto de vista, os rapazes teriam sempre mais irmãs do que as raparigas.

Isto parece paradoxal. Analisemos a situação mais atentamente. Se uma família com n filhos tiver m meninas e $n-m$ rapazes, os rapazes terão um total de $m(n-m)$ irmãs entre todos, ao passo que as raparigas terão um total de $m(m-1)$ irmãs. Estes números só podem ser iguais se $m = \frac{1}{2}(n+1)$. Isto nunca pode ser verdade se n for um número par pois, nesse caso, m seria uma fração.

* Nas áreas rurais é permitido ter um segundo filho depois de um intervalo de 3 anos se o primeiro filho for deficiente ou do sexo feminino.

Este enigma foi criado porque uma família com filhos pode compreender várias combinações de ambos os sexos. Uma família com 3 filhos pode ter 3 rapazes, 3 meninas, 2 rapazes e 1 menina, ou 2 meninas e 1 rapaz. Se supusermos que existe uma probabilidade igual de $\frac{1}{2}$ de um recém-nascido ser rapaz ou rapariga (o que não é propriamente verdade), uma família com n filhos pode ser constituída de 2^n maneiras diferentes. O número de diferentes composições de uma família com n filhos e m meninas é denotado* por nC_m e cada um dos rapazes terá $m(n-m)$ irmãs. Considerando as 2^n formas diferentes pelas quais a família pode ser dividida em irmãos e irmãs, deveríamos estar a perguntar qual é o número médio de irmãs que os rapazes nas famílias com n filhos terão. Este número é a soma de todas as respostas para o número de irmãs que podem ter para todos os valores possíveis do número $m = 0, 1, 2, \dots, n$ dividido pelo número total, 2^n . Ou seja:

$$r_n = 2^{-n} \sum_m {}^nC_m \times m(n-m)$$

De forma semelhante, o número médio de irmãs que as raparigas em famílias com n filhos terão é:

$$m_n = 2^{-n} \sum_m {}^nC_m \times m(m-1)$$

As soluções para estas fórmulas são muito mais simples do que as fórmulas podem levar a crer. Curiosamente, o número médio de irmãs para os rapazes e raparigas é igual e $r_n = m_n = \frac{1}{4}n(n-1)$. É de notar que, por se tratar de uma média, isto não significa que uma determinada família tenha de ter o comportamento da média. Quando $n = 3$, o número médio de irmãs é 1,5 (que nenhuma família pode ter). Quando $n = 4$, o número médio é 3. Quando n se torna grande, o número médio aproxima-se cada vez mais do quadrado de $n/2$. Apresento abaixo a tabela representativa das 8 possibilidades para famílias com 3 filhos:

* nC_m é a representação abreviada de $n!/(m!(n-m)!)$ e corresponde ao número de formas de escolher m resultados a partir de um total de n possibilidades.

Constituição da família com 3 filhos	Número de formas pelas quais a família pode ser composta	Número de irmãs para os rapazes	Número de irmãs para as meninas
3 rapazes	1	0	0
2 rapazes + 1 menina	3	2	0
2 meninas + 1 rapaz	3	2	2
3 meninas	1	0	6

Podemos constatar que o número total de irmãs para os rapazes corresponde às somas da segunda e terceira filas, $3 \times 2 + 3 \times 2 = 12$ e que o número total de irmãs para as meninas corresponde à soma da terceira e quarta filas; também é igual a $12 = 3 \times 2 + 1 \times 6$. Uma vez que esta família pode ter 8 combinações de ambos os sexos, o número médio de irmãs para as meninas e para os rapazes é igual a $12/8 = 1,5$, conforme previsto pela nossa fórmula $\frac{1}{4} \times n \times (n-1)$ para o caso de $n = 3$, quando existem 3 crianças na família.

69 | JOGAR LIMPO COM UMA MOEDA VICIADA

E, o que é surpreendente, Cambridge ganhou o lançamento.

Harry Carpenter

Há alturas em que precisamos de lançar uma moeda que sabemos que não está viciada para escolhermos entre duas opções de uma forma não tendenciosa. No início de muitos eventos desportivos, o árbitro lança uma moeda ao ar e pede aos participantes que escolham “face” ou “verso”. Pode ser feito um jogo de azar, marcando sequências de lançamentos de uma moeda. Podemos usar mais do que uma moeda simultaneamente, de modo a criar um grande número de resultados possíveis. Agora suponhamos que a única moeda que tem disponível é uma moeda viciada: que não tem probabilidades iguais (de $\frac{1}{2}$) de dar “face” ou “verso”. Ou talvez suspeite de que a moeda que o seu adversário atenciosamente forneceu está viciada. Há alguma coisa que possa fazer para se certificar de que o lançamento de uma moeda viciada produz dois resultados igualmente prováveis e justos?

Suponhamos que lança a moeda ao ar duas vezes e ignora os casos em que ambos os lançamentos produzem o mesmo resultado – ou seja, que repete o lançamento se obtiver uma sequência de “face-face” (FF) ou “verso-verso” (VV). Daqui podem resultar dois pares de resultados: “face” seguido de “verso” (FV) ou “verso” seguido de “face” (VF). Se a probabilidade de a moeda viciada ter como resultado “face” for p , a probabilidade de obter como resultado “verso” será $1-p$, sendo a probabilidade de obter a sequência FV $p(1-p)$ e a de obter VF $(1-p)p$. Estas duas probabilidades

* À semelhança do Capítulo 24, neste texto optou-se pelas designações “face” e “verso” em vez de “cara” e “coroa” para se poderem obter iniciais diferentes. (N. T.)

são iguais, independentemente da probabilidade p das moedas viciadas. Para garantirmos que o jogo é justo, basta-nos definir FACE como a sequência FV e VERSO como a sequência VF, e a probabilidade de o resultado ser VERSO passa a ser igual à do resultado ser FACE. E não precisa de conhecer a tendência, p , da moeda*.

* Este truque foi concebido pelo grande matemático, físico e pioneiro da informática John von Neumann. Foi muito usado na construção dos algoritmos informáticos. Uma das questões que lhe foram subsequentemente colocadas foi se haveria uma forma mais eficiente de definir os novos estados de FACE ou VERSO. A forma que aqui apresentámos faz-nos desperdiçar "tempo" a ter de descartar todas as sequências FF e VV obtridas.

70 | AS MARAVILHAS DA TAUTOLOGIA

Uma boa forma de compreender os acontecimentos do mundo moderno é supor o contrário do que o Lord Rees-Mogg afirma.

Richard Ingrams

“**T**autologia” é uma palavra com más vibrações. Sugere ausência de sentido e o meu dicionário define-a como “repetição desnecessária de uma ideia, frase ou palavra”. É uma afirmação que é verdadeira em todas as circunstâncias: todos os cães ruivos são cães. Mas seria errado pensar que as tautologias não são úteis. De certa forma, podem ser a única forma de obter conhecimento. Apresento em seguida uma situação em que a sua vida depende de uma tautologia.

Imagine que está preso numa cela com duas portas – uma porta vermelha e uma porta preta. Uma destas portas – a vermelha – conduz à morte certa e a outra – a preta – conduz à segurança, mas o leitor não sabe qual das portas conduz a qual cenário. Cada uma das portas tem um telefone ao lado que pode usar para falar com um conselheiro que lhe dirá que porta deve escolher para sair da cela em segurança. O problema é que um dos conselheiros diz sempre a verdade e o outro mente sempre, mas o leitor não sabe com qual dos dois está a falar. Só pode fazer uma pergunta. O que deve perguntar?

Consideremos a pergunta mais simples que pode fazer: “Que porta devo escolher?” O conselheiro que diz sempre a verdade dir-lhe-á que escolha a porta preta e o que mente sempre dir-lhe-á que opte pela porta vermelha. Como o leitor não sabe qual dos dois conselheiros está a dizer a verdade, esta informação não é útil. Teria igual probabilidade de acertar se escolhesse aleatoriamente uma das portas. “Que porta devo escolher?” não é, portanto, uma tautologia nesta situação. É uma pergunta que admite respostas diferentes.

Suponhamos que perguntava: “Que porta é que o outro conselheiro me recomendaria que escolhesse?” Esta situação é mais interessante. O conselheiro que diz sempre a verdade sabe que o conselheiro que mente recomendaria sempre que escolhesse a porta vermelha – que conduz à morte certa – e assim, o conselheiro honesto dir-lhe-á que o outro conselheiro lhe recomendaria a porta vermelha. O conselheiro mentiroso sabe que o conselheiro honesto lhe recomendaria que escolhesse a porta preta para sair da cela em segurança e vai tentar enganá-lo, dizendo que ele lhe recomendaria a porta vermelha.

Acaba de fazer uma descoberta que lhe vai salvar a vida. Independentemente de quem esteja a responder à sua pergunta, recomendam-lhe que escolha a porta vermelha. Encontrou uma pergunta que é uma tautologia – resulta sempre na mesma resposta – e que é a sua salvação. Assim, a estratégia para sair da cela em segurança é perguntar “Que porta é que o outro conselheiro me recomendaria que escolhesse?”, tomar nota da resposta (vermelha) e escolher a outra porta (preta) para alcançar a segurança.

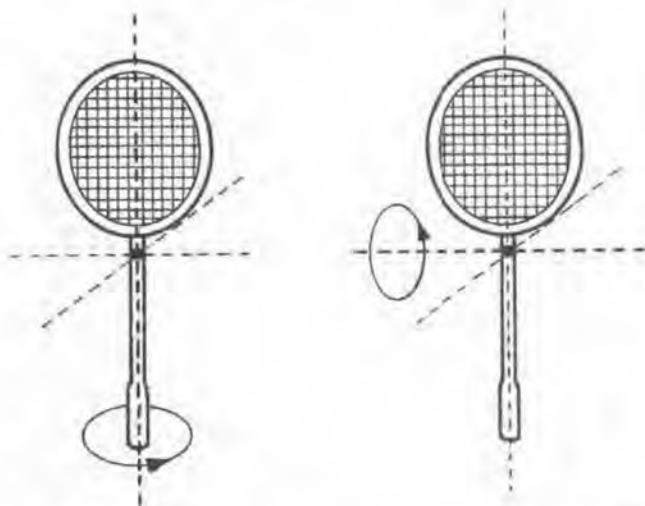
A velocidade nunca matou ninguém; é parar de repente que nos mata.

Jeremy Clarkson

Algumas coisas são mais difíceis de mover do que outras. A maioria das pessoas pensa que a única coisa que conta neste tipo de problemas é o peso. Quanto maior for a carga, mais difícil será movê-la. Mas tentar mover diferentes tipos de cargas, descobrirá que a concentração da massa tem um papel importante. Quanto mais concentrada estiver a massa, mais fácil de mover será e mais depressa oscilará (lembre-se do que aprendemos no Capítulo 2). Consideremos uma patinadora no gelo a iniciar uma pirueta. Começará por estender os braços para os lados e depois aproximá-los-á do corpo. Isto resulta num rápido aumento da velocidade de rotação. À medida que a massa da patinadora fica mais concentrada junto do seu centro, ela mover-se-á mais depressa. Por outro lado, se analisarmos as formas das vigas usadas para construir edifícios robustos, estas têm uma secção transversal em forma de H, que distribui a massa mais longe do centro da viga, tornando mais difícil movê-la ou deformá-la quando posta sob tensão.

A resistência à deslocação chama-se “inércia” no uso corrente e é determinada pela massa total de um objeto e pela sua distribuição, que será determinada pela forma do objeto. Se continuarmos a pensar na rotação, um exemplo interessante será o de um objeto simples como uma raquete de ténis. Esta tem uma forma invulgar e pode ser girada de três formas distintas. Podemos pousá-la no chão e girá-la em torno do seu centro. Podemos apoiá-la no ponto superior e girá-la pela pega. E podemos agarrar nela pela pega e atirá-la ao ar, de modo a que dê uma cambalhota, e volte a descer para ser novamente agarrada pela pega. Existem três formas de a girar, porque existem três direções no espaço, cada uma em ângulos

retos em relação às outras duas, e a raquete pode ser girada em torno de um eixo em qualquer uma destas três direções. A raquete comporta-se de forma diferente quando é girada em torno de cada um dos diferentes eixos por ter uma inércia diferente ao movimento em torno de cada um deles. Apresento abaixo dois exemplos:



Estes movimentos diferentes têm uma propriedade notável, que é revelada pelos três lançamentos diferentes da raquete. O movimento em torno dos eixos em volta dos quais a inércia é maior ou menor é simples. Quando a raquete está deitada no chão ou apoiada sobre a parte superior, não faz nada especialmente invulgar. Mas quando a giramos em torno do eixo intermédio, em volta do qual a inércia tem um valor entre o maior e o menor (apresentados à direita), acontece algo invulgar. Pegue na raquete pela pega, com a face voltada para cima, como se fosse uma frigideira. Marque a face que está voltada para cima com giz. Lance a raquete ao ar de modo a que dê uma volta de 360 graus sobre si própria e volte a agarrar nela pela pega. A face que marcou com o giz estará agora virada para baixo.

A regra de ouro diz que uma volta em torno do eixo com a inércia intermédia é instável. O mais ligeiro desvio da linha central exata fá-la virar-se. Às vezes, isto é uma coisa boa. Se for uma ginasta a fazer uma cambalhota na trave, o exercício parecerá mais difícil (e garantir-lhe-á uma melhor pontuação) se fizer também uma rotação. Mas essa rotação pode acontecer automaticamente por causa desta instabilidade.

Um exemplo mais grave desta instabilidade surgiu há alguns anos na Estação Espacial Internacional, depois de esta ter sido atingida durante uma aterragem mal calculada de uma nave de abastecimento russa. A Estação sofreu danos e começou a girar lentamente. Ainda havia gases no sistema de retrofoguetes, que podiam ser acionados para abrandar esta rotação e devolver a Estação ao seu estado habitual de equilíbrio. No entanto, punha-se o problema de como disparar os foguetes. Em que direção deveriam mover a Estação para contrariar a rotação existente? O astronauta britânico Michael Foales teve de resolver este problema enquanto estava preso na parte não danificada da Estação, usando o seu computador portátil e uma ligação à Terra. Os aspetos mais importantes a averiguar eram as três inércias da Estação Espacial relativamente às rotações em torno dos seus três eixos. Se os foguetes fossem acionados de forma errada, podiam lançar a Estação numa rotação em torno do seu eixo de inércia intermédia. O resultado teria sido uma verdadeira catástrofe. A instabilidade que virou a nossa raquete de ténis não teve um efeito prejudicial sobre a raquete, mas, se virasse a Estação Espacial, esta desintegrar-se-ia, matando os astronautas, espalhando no espaço um quarto de um milhão de quilos de detritos potencialmente mortais e causando perdas financeiras astronómicas. A NASA não conhecia as três inércias da Estação Espacial – ninguém pensou que esse conhecimento viesse a ser necessário – e, assim, Foales teve de as calcular a partir das plantas da Estação Espacial e depois determinar a forma como ela reagiria aos foguetes disparados em diferentes direções para corrigirem a rotação produzida pelo acidente. Felizmente ele conhecia a instabilidade da rotação em torno do eixo intermédio e fez os cálculos certos. A rotação perigosa foi corrigida e os astronautas foram salvos. A matemática, por vezes, pode ser uma questão de vida ou morte.

Em viagem: aprendi que precisamos de quatro vezes mais água, do dobro do dinheiro e de metade da roupa que pensamos que vamos precisar quando partimos.

Gavin Esler

Certa vez mostraram a um jovem um frasco de vidro vazio com uma tampa de rosca. A seguir, deram-lhe uma caixa de bolas de tênis e pediram-lhe que enchesse o frasco. Ele deitou algumas bolas de tênis lá para dentro, depois abanou o frasco para tentar enfiar lá mais uma bola antes de colocar a tampa. Perguntaram-lhe: "O frasco está cheio?" E o rapaz respondeu: "Sim, está cheio." Depois, deram-lhe uma caixa de berlines e pediram-lhe que tentasse pôr alguns dentro do frasco. O rapaz abriu a tampa e descobriu que conseguia introduzir uma boa quantidade de berlines entre as bolas de tênis. Abanando ocasionalmente o frasco, conseguia encaixar ainda mais alguns berlines nos espaços vazios. Quando já não conseguia colocar mais nenhum berline no frasco, anunciou que este estava cheio. O mentor pegou então num saco de areia e pediu ao rapaz que enchesse o frasco. Ele voltou a desenroscar a tampa e deitou a areia para dentro do frasco. Desta vez não precisou de deslocar muito o conteúdo do frasco, bastava-lhe apenas abaná-lo cuidadosamente de vez em quando para se certificar de que a areia estava a escoar para todos os espaços vazios entre as bolas de tênis e os berlines. Finalmente, quando já não conseguia deitar mais areia para dentro do frasco, voltou a colocar a tampa. Agora o frasco estava verdadeiramente cheio!

Podemos aprender várias lições com esta história. Se tivessem começado por dar a areia ao rapaz e por lhe pedir que a usasse para encher o frasco, não teria sobrado espaço para os berlines ou para as bolas de tênis. Precisamos de começar pelas coisas maiores se queremos arranjar

espaço para elas. O mesmo é aplicável a problemas de organização mais familiares. Se precisar de colocar várias caixas numa carrinha, é útil pensar na forma como deve colocá-las de modo a garantir que as consegue pôr lá todas. A nossa pequena história mostra-nos que devemos começar pelos objetos maiores e depois avançar para os segundos maiores e por aí fora, deixando os objetos mais pequenos para o fim.

As formas dos objetos que estamos a tentar arrumar também são, obviamente, importantes. Frequentemente, estes objetos têm todos o mesmo tamanho. Se o leitor for um fabricante de doces ou de outros alimentos pequenos, poderá querer saber qual é a forma ideal para conseguir incluir o máximo deles num frasco ou noutro tipo de embalagem. Durante algum tempo, pensou-se que a melhor solução era dar-lhes a forma de pequenas esferas. Muitas esferas pequenas pareciam deixar o menor espaço vazio possível entre si. Curiosamente, veio a descobrir-se que esta não era a forma ideal. Se os doces fossem feitos na forma de pequenos elipsoides, como bolas de rãguebi em miniatura ou amêndoas, podiam preencher ainda mais espaço. Desta forma, os Smarties e os M&Ms enchem um volume de forma mais eficiente do que qualquer conjunto de esferas. Se os elipsoides tiverem os eixos curto e longo na proporção de 1 para 2, deixarão apenas 32% do espaço desocupado, em comparação com os 36% deixados por doces de forma esférica. Este facto aparentemente trivial tem muitas consequências importantes para a eficiência comercial e para a produção industrial, reduzindo o desperdício, os custos de expedição e evitando o uso desnecessário de embalagens.

73 | FAZER NOVAMENTE AS MALAS

Tenho as malas feitas; estou pronto para partir.

John Denver

O nosso pequeno problema de embalagem representado pelo exemplo do frasco do capítulo anterior era um problema simples. Começamos com os objetos maiores e depois fomos usando objetos progressivamente mais pequenos até chegarmos aos mais pequenos de todos. Na prática, o problema pode ser mais complexo. Podemos ter muitos recipientes para encher com produtos comprados, todos eles com tamanhos diferentes. Como é que distribuimos os artigos de diferentes tamanhos pelas embalagens de modo a usar o menor número de embalagens? “Embalar” pode não se referir apenas a arrumar num espaço determinado; pode significar organizar no tempo. Suponhamos que o leitor é gerente de um grande centro de cópias que faz reproduções de diferentes documentos com diferentes dimensões para os seus clientes. Como é que distribui os diferentes trabalhos pelas máquinas de modo a minimizar o número de máquinas que precisa de utilizar para fazer o trabalho de um dia?

Tudo isto são versões de um problema que os computadores demoram muito tempo a resolver quando o número de artigos a embalar e o número de “caixas” a usar se torna grande.

Imaginemos que pode usar caixas que permitem armazenar um máximo de 10 embalagens e que lhe são dadas 25 embalagens de tamanhos distintos para guardar nestas caixas de modo a usar o menor número possível de caixas da forma mais eficiente possível. As dimensões das embalagens individuais são as seguintes:

6, 6, 5, 5, 5, 4, 3, 2, 2, 3, 7, 6, 5, 4, 3, 2, 2, 4, 4, 5, 8, 2, 7, 1

Antes de mais, imaginemos que estas embalagens chegam num tapete rolante, pelo que não pode separá-las como um grupo: só pode empilhá-las uma a uma à medida que chegam. A estratégia mais simples a seguir é ir colocando as embalagens na primeira caixa até não caberem mais e depois começar a usar uma nova caixa. Não pode voltar atrás e preencher os espaços vazios nas caixas anteriores porque estas são levadas dali para fora. A estratégia é, por vezes, designada por "Encaixe Seguinte". A embalagem de tamanho 6 vai para a primeira caixa. A segunda embalagem, de tamanho 6, já não cabe na mesma caixa, portanto abre uma nova. A embalagem de tamanho 5 que se segue não caberá nesta caixa, portanto abrirá uma terceira. Quando acrescentar a segunda embalagem de tamanho 5, a caixa ficará cheia e as duas embalagens de tamanho 5 que se seguem encherão a caixa seguinte e por aí fora. Eis o resultado que obtemos se seguirmos este método do Encaixe Seguinte para as embalagens encontradas pela ordem indicada acima:

[6], [6], [5,5], [5,5], [4,3,2], [2,3], [7], [6], [5,4], [3,2,2], [4,4], [5], [8,2], [7,1]

Usámos 14 caixas e apenas 3 delas estão cheias (as duas caixas [5,5] e a caixa [8,2]). A quantidade total de espaço não preenchido nas caixas que não estão cheias até ao limite da sua capacidade é $4+4+1+5+3+4+1+3+2+5+2=34$.

A grande quantidade de espaço desperdiçado foi causada pelo facto de não podermos voltar atrás e preencher os espaços vazios nas caixas anteriores. Conseguiria obter um resultado melhor se pudesse usar uma estratégia que lhe permitisse colocar as embalagens na primeira caixa disponível que tivesse espaço? Esta estratégia é, por vezes, chamada "Primeiro Encaixe". Usando a estratégia do Primeiro Encaixe, começamos como anteriormente, com duas embalagens de tamanho 6 em duas caixas distintas e depois enchemos duas caixas com dois pares de embalagens de tamanho 5. No entanto, a embalagem seguinte é de tamanho 4, que podemos colocar na primeira caixa, junto com a embalagem de tamanho 6; em seguida, podemos colocar duas embalagens de tamanho 2 e uma embalagem de tamanho 3 na quinta caixa, e por aí em diante, até

chegarmos à última embalagem, de tamanho 1, que colocamos na segunda caixa, enchendo-a. É esta a distribuição final das embalagens:

$$[6,4], [6,3,1], [5,5], [5,5], [2,2,3,3], [7,2], [6,4], [5,2,2], [4,4], [5], [8], [7]$$

A estratégia do Primeiro Encaixe teve um resultado muito melhor do que a do Encaixe Seguinte. Só usámos 12 caixas e a quantidade de espaço desperdiçado foi reduzida para $1+1+2+5+2+3=14$. Conseguimos encher completamente 6 das caixas.

Agora começamos a perceber como ainda podemos alcançar melhores resultados nesta atividade da embalagem. O espaço desperdiçado surge quando temos uma embalagem grande no fim da lista. Os espaços deixados nas primeiras caixas são pequenos quando chegamos ao fim, e somos obrigados a usar uma nova caixa para cada embalagem nova. É óbvio que poderíamos ter resultados melhores se pudéssemos organizar as embalagens por ordem decrescente de tamanhos. Esta opção pode não ser sempre possível, por exemplo quando estamos a organizar bagagens que são descarregadas de um avião por um tapete rolante, mas vejamos até que ponto pode ser útil nos casos em que é.

Se organizarmos as nossas embalagens por ordem decrescente de tamanhos, obteremos a seguinte lista:

$$8,7,7,6,6,6,5,5,5,5,5,5,4,4,4,4,3,3,3,2,2,2,2,1$$

Agora experimentemos usar a nossa antiga estratégia do Encaixe Seguinte depois de feita esta seleção – chamemos-lhe “Encaixe Seguinte Ordenado”. As primeiras seis embalagens vão em caixas separadas, depois enchemos três caixas com pares de embalagens de tamanho 5, etc. O resultado final tem este aspeto:

$$[8], [7], [7], [6], [6], [6], [5,5], [5,5], [5,5], [4,4], [4,4], [3,3,3], [2,2,2,2,2], [1]$$

Tivemos azar com a última embalagem! Tivemos de usar uma caixa nova para colocar apenas uma embalagem de tamanho 1. Usando a estratégia do Encaixe Seguinte Ordenado, acabámos a precisar novamente de 14 caixas – tal como aconteceu antes da seleção – e voltámos

a ter um desperdício de 34. É exatamente o mesmo resultado que obtivemos com o Encaixe Seguinte não ordenado. Mas se não tivéssemos incluído aquela última embalagem de tamanho 1, teríamos continuado a precisar de 14 caixas para o Encaixe Seguinte, mas apenas de 13 caixas para o Encaixe Seguinte Ordenado.

Finalmente, vejamos o que acontece se usarmos uma estratégia do "Primeiro Encaixe Ordenado". Mais uma vez, as embalagens de tamanho 6 são colocadas em caixas separadas, depois as seis embalagens de tamanho 5 enchem mais três caixas, mas depois a seleção acontece espontaneamente. Três das embalagens de tamanho 4 enchem as caixas que contêm as embalagens de tamanho 6, ao passo que a outra embalagem de tamanho 4 abre uma nova caixa. As embalagens restantes enchem facilmente os espaços vazios, deixando apenas a última caixa incompleta:

$$[8,2],[7,3],[7,3],[6,4],[6,4],[6,4],[5,5],[5,5],[5,5],[4,3,2,1],[2,2,2]$$

Usámos 11 caixas e o único espaço desperdiçado foi de valor 4 na última caixa. Este resultado é muito superior ao que obtivemos com as nossas outras estratégias e podemos perguntar-nos se será o melhor resultado possível. Pode haver outra estratégia de embalagem que use ainda menos do que 11 caixas para esta quantidade? É fácil perceber que não é possível. O tamanho total das embalagens é $1 \times 8 + 2 \times 7 + 3 \times 6 + \dots + 5 \times 2 + 1 \times 1 = 106$. Uma vez que cada caixa só pode conter embalagens com um tamanho total de valor 10, a totalidade das embalagens requer a utilização de, pelo menos, $106/10 = 10,6$ caixas. Desta forma, nunca podemos usar menos do que 11 caixas e nunca poderemos desperdiçar menos do que 4 espaços.

Assim sendo, encontrámos a melhor solução possível usando o método do Primeiro Encaixe Ordenado. Se voltarmos a analisar o problema simples que apresentámos no capítulo anterior, o de encaixar objetos de três tamanhos num frasco, percebemos que estávamos a usar a estratégia do Primeiro Encaixe Ordenado, porque colocámos primeiro os objetos maiores e só depois os mais pequenos. Infelizmente, os problemas de organização não são sempre assim tão fáceis. De um modo geral, não há um método rápido que permita a um computador encontrar a melhor solução de organização de uma determinada seleção de embalagens no menor

número de caixas. À medida que os tamanhos das caixas e a variedade de dimensões das embalagens aumentam, o problema torna-se muito difícil e, se o número de embalagens crescer o suficiente e se os seus tamanhos forem suficientemente diversificados, acabará por derrotar qualquer computador que tente encontrar a melhor distribuição das embalagens pelas caixas num determinado período de tempo. Mesmo no caso deste problema, há outras considerações que podem levar-nos a concluir que o Primeiro Encaixe Ordenado é apenas a segunda melhor solução. A seleção das embalagens, da qual depende a eficiência deste método, consome tempo. Se o tempo despendido a colocar as embalagens em caixas também for tido em conta, simplesmente usar menos caixas pode não ser a solução mais económica.

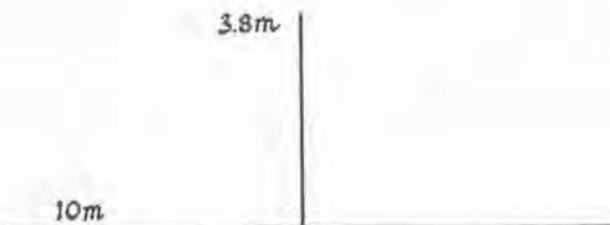
Tigre, tigre, que flamejas nas florestas da noite.

William Blake, "O Tigre"

Há pouco tempo, houve um trágico acidente no zoo de São Francisco, quando Tatiana, um pequeno (!) tigre siberiano com 135 quilos saltou o muro que encerrava o seu habitat, matou um visitante e deixou outros dois gravemente feridos. Os jornais citaram os responsáveis do zoo que disseram que estavam espantados com as evidências de que o tigre tinha conseguido saltar por cima do muro alto que rodeava o seu habitat: "Tem de ter saltado. Mas como é que conseguiu saltar uma altura tão alta é que me parece espantoso", disse o diretor do zoo, Manuel Mollinedo. Embora se tenha afirmado inicialmente que o muro que rodeava o habitat dos tigres tinha 5,5 metros de altura, mais tarde veio a descobrir-se que a sua altura não passava dos 3,8 metros, uma altura muito inferior aos 5 metros recomendados pela Associação Americana de Zoos e Aquários para garantir a segurança. Mas devemos sentir-nos seguros com algum destes números? A que altura consegue um tigre saltar?

O muro estava rodeado por um fosso seco com 10 metros de largura, pelo que um tigre cativo se deparava com o desafio de saltar a uma altura de 3,8 metros, partindo de uma corrida de balanço horizontal de pelo menos 10 metros até ao muro. Em distâncias curtas e em terreno plano, um tigre consegue atingir velocidades de mais de 22 metros por segundo. Com uma corrida de balanço de 5 metros, pode alcançar facilmente uma velocidade de 14 metros por segundo no momento do salto.

O problema é exatamente igual ao do lançamento de um projétil. Segue um percurso parabólico até à altura máxima e depois desce. A velocidade



mínima de lançamento V que alcançará uma altura vertical, a , a partir do ponto de lançamento a uma distância x do muro é determinada pela fórmula:

$$V^2 = g(a + \sqrt{a^2 + x^2})$$

Em que $g = 10 \text{ m/s}^2$ corresponde à aceleração produzida pela gravidade da Terra. Saliento algumas características desta equação que provam que faz sentido: se tornarmos a gravidade mais forte (se o valor g se tornar maior), será mais difícil saltar e a velocidade mínima de lançamento, V_m , tem de ser maior. Da mesma forma, se a altura do muro, a , aumentar ou se a distância a que se encontra no momento do lançamento, x , aumentar, é necessária uma maior velocidade de lançamento para conseguir transpor o muro.

Consideremos a configuração do habitat dos tigres no zoo de São Francisco, conforme já apresentada. O muro tem 3,8 metros de altura, mas o tigre é corpulento e o seu centro tem de se elevar a cerca de 4,3 metros para conseguir transpor o muro, uma vez que os tigres siberianos têm cerca de um metro de altura na região dos ombros. (Vamos ignorar a possibilidade de ele conseguir agarrar-se ao muro para ganhar impulso – que é, apesar de tudo, bastante provável.) O resultado é $V^2 = 9,8(4,3 + \sqrt{(18,5 + 100)}) = 148,97 \text{ (m/s)}^2$, pelo que V é igual a 12,2 m/s.

Este valor está dentro da velocidade de lançamento possível para um tigre, pelo que ele deveria mesmo ter conseguido saltar o muro. Se o muro fosse aumentado para os 5,5 metros de altura, o tigre já teria de elevar o seu centro de gravidade a 6 metros para conseguir passar o muro e, para tal, precisaria de alcançar uma velocidade de lançamento de 13,2 metros por segundo. Como disse o diretor: "Obviamente, agora que algo aconteceu, teremos de reconsiderar a altura real do muro."

* Nestes problemas, considera-se o projétil como uma massa de tamanho negligenciável localizada no seu centro (também chamada "massa pontual"). Claro que o tigre tem um tamanho considerável que não é, de forma alguma, um ponto. Ainda assim, vamos ignorar este facto e considerar o tigre como se tivesse toda a massa localizada no seu centro.

Então, o Eriope juntou os cinco dedos... e encostou-os ao Leopardo, e onde os dedos tocaram apareceram cinco pequenas manchas pretas, todas muito juntas... Ocasionalmente, os dedos deslizavam e as manchas ficavam um pouco esbatidas, mas se olhares com atenção para qualquer Leopardo, notarás que existem sempre cinco manchas juntas – as das pontas dos cinco dedos.

Rudyard Kipling, "O leopardo e as suas manchas", in *Histórias Assim Mesmo*

Os padrões no pelo dos animais, especialmente dos grandes felinos, são uma das imagens mais espetaculares que podemos encontrar no mundo natural. Estes padrões não são de forma alguma aleatórios, nem tão pouco são determinados exclusivamente pela necessidade de camuflagem. Os ativadores e inibidores que encorajam ou bloqueiam a presença de pigmentos específicos fluem no embrião do animal de acordo com uma lei simples que dita a forma como a sua concentração em pontos específicos depende da produção desse pigmento por reações químicas e a velocidade a que se espalha pela pele. O resultado é uma difusão de sinais semelhante a uma onda, sinais esses que ativam ou suprimem diferentes pigmentos. Os efeitos resultantes dependem de várias coisas, como o tamanho e a forma do animal, e o comprimento das ondas de padrões. Se observar uma área da superfície da pele, reparará que os picos e depressões destas ondas criam uma rede regular de montes e vales de cores diferentes. O surgimento de um pico ocorre como consequência da tendência inibidora, resultando assim numa risca ou mancha pronunciada contra um fundo de cor contrastante. Se se der a concentração máxima possível num determinado lugar, o aumento dessa concentração acabará por ter de se espalhar e as manchas acabarão por se fundir, transformando-se em manchas maiores ou riscas.

O tamanho do animal é importante. Um animal muito pequeno não terá espaço para permitir que muitos altos e baixos da onda ativadora dos

pigmentos se espalhem por todo o seu corpo, pelo que acaba por ser de uma só cor ou, na melhor das hipóteses, malhado como um hamster. Quando o animal é muito grande, como um elefante, o número de altos e baixos das ondas é tão grande que o efeito geral é monocromático. No tamanho intermédio, entre o muito grande e o muito pequeno, há espaço para uma grande variedade, tanto de um animal para outro como na superfície do corpo de um único animal. A chita, por exemplo, tem o corpo às manchas, mas a cauda às riscas. As ondas criam picos e depressões distintos à medida que se espalham pelo corpo grande e moderadamente cilíndrico da chita mas, quando se espalham para a cauda estreita e cilíndrica, estão muito mais juntas e fundem-se para criar a aparência de riscas. Esta tendência origina um "teorema" matemático muito interessante que provém do comportamento das ondas de concentração de cor nos corpos dos animais: os animais com o pelo às manchas podem ter caudas às riscas, mas os animais com o pelo às riscas não podem ter caudas às manchas.

O futuro pertence às multidões.

Don DeLillo, *Mao II*

Se já alguma vez estive no meio de uma grande multidão, como num evento desportivo, num concerto de música *pop* ou numa manifestação, pode ter sentido ou testemunhado algumas características estranhas do comportamento coletivo das pessoas. A multidão não está a ser organizada como um todo. Todas as pessoas reagem ao que está a acontecer ao seu lado, mas a multidão pode, ainda assim, modificar subitamente o seu comportamento numa área muito vasta, com resultados desastrosos. Uma procissão a um ritmo lento pode transformar-se numa debandada motivada pelo pânico, com pessoas a tentarem correr em todas as direções. É importante compreender esta dinâmica. Se houver um incêndio ou uma explosão perto de uma grande multidão, como irão comportar-se as pessoas? Que tipo de rotas de fuga e saídas gerais devem ser planeadas em grandes estádios? Como devem ser organizadas as peregrinações religiosas de milhões de crentes a Meca de modo a evitar que se repitam as mortes de centenas de peregrinos que já se verificaram no passado, quando o pânico gera uma debandada humana como reação à sobrelotação do espaço?

Um dos factos interessantes que informam os estudos recentes sobre o comportamento das multidões e o seu controlo é a analogia entre o fluxo de uma multidão e o fluxo de um líquido. A princípio poderíamos pensar que a tentativa de compreender os comportamentos de multidões compostas por diferentes tipos de pessoas, todos eles com diferentes reações potenciais a uma determinada situação e com diferentes idades e graus de compreensão dessa mesma situação, seria uma tarefa impossível, mas surpreendentemente não é verdade. As pessoas são mais parecidas do que imaginamos. Simples escolhas de local podem resultar em ordem genera-

lizada num contexto de sobrelotação do espaço. Ao chegar a um dos grandes terminais do caminho de ferro em Londres e dirigir-se para o metropolitano, descobre que as pessoas que estão a descer escolhem a escada da esquerda (ou da direita), ao passo que as que estão a subir escolhem a outra. Ao longo dos corredores que conduzem às barreiras que controlam as entradas e saídas, a multidão organiza-se em duas correntes distintas a avançar em direções opostas. Ninguém planeou as coisas desta forma nem foram afixadas placas a exigir que assim fosse: esta situação derivou de um comportamento individual formado a partir do que as pessoas viam ao seu redor. Isto significa que elas agem em resposta à forma como as outras pessoas se movem perto de si e em relação à concentração de pessoas no mesmo espaço. As reações ao segundo fator dependem muito da pessoa em questão. Se se tratar de um homem de negócios japonês habituado a viajar no sistema de transportes públicos de Tóquio em hora de ponta, reagirá de forma muito diferente a uma grande multidão de pessoas à sua volta do que um turista escocês ou um grupo de estudantes de Roma. Se estiver a acompanhar um familiar mais jovem ou mais velho, movimentar-se-á de uma forma específica, muito próximo deles e sempre atento ao seu paradeiro. Todas estas variáveis podem ser ensinadas a computadores que poderão simular o que acontecerá a multidões reunidas em espaços diferentes e a forma como reagirão ao desenvolvimento de novas tensões.

O comportamento das multidões parece compreender três fases, tal como o fluxo de um líquido. Quando a acumulação de pessoas não é demasiado grande e o movimento da multidão segue firmemente numa direção – como a multidão a abandonar o estádio de Wembley na direção da estação de metro de Wembley Park depois de um jogo de futebol – comporta-se como um fluxo muito suave de um líquido. Todas as pessoas começam a mover-se mais ou menos à mesma velocidade em vez de pararem e arrancarem constantemente.

No entanto, se a densidade de pessoas numa multidão aumentar significativamente, as pessoas começarão a empurrar-se e começará a dar-se um movimento em diferentes direções. O movimento geral torna-se mais irregular, com um comportamento mais do tipo “para-arranca” em vez de se comportar como uma sucessão de ondas. O aumento gradual da densidade dos corpos reduzirá a velocidade a que podem avançar e algumas pessoas tentarão deslocar-se para os lados se sentirem que podem

avançar mais depressa assim. É exatamente a mesma psicologia dos condutores que mudam constantemente de faixa num engarrafamento denso e em progressão lenta. Em ambos os casos, criam-se ondas que alastram pelo engarrafamento, o que leva algumas pessoas a abrandarem e outras a deslocarem-se para o lado para deixarem outro condutor entrar na sua faixa. Uma sucessão dessas ondas alastra-se pela multidão. Na sua essência, estas ondas não são necessariamente perigosas, mas indicam a possibilidade de algo mais perigoso acontecer de repente.

Uma concentração muito grande de pessoas numa multidão acabará por fazê-las comportarem-se de uma forma muito mais caótica, como o fluxo de um líquido que fique subitamente agitado, à medida que as pessoas começam a tentar mover-se em várias direções, para encontrarem espaço livre. Empurram as pessoas à sua volta e as suas tentativas de criar espaço pessoal tornam-se mais vigorosas. Isto aumenta o risco de quedas e de as pessoas ficarem tão apertadas que têm dificuldade em respirar ou de as crianças se perderem dos pais. Estes efeitos podem começar em pontos diferentes de uma grande multidão e os seus efeitos podem espalhar-se muito rapidamente. Cria-se um efeito bola de neve e a situação descontrola-se. As pessoas caídas transformam-se em obstáculos, por cima dos quais caem outras pessoas. Uma pessoa com claustrofobia entraria em pânico muito rapidamente e reagiria de forma ainda mais violenta às pessoas à sua volta. A menos que ocorra algum tipo de intervenção organizada que separe a multidão em várias partes e reduza a densidade de pessoas naquele espaço, o desastre é iminente.

A transição de um fluxo suave de peões para um movimento irregular e para o caos da multidão pode demorar entre poucos minutos e meia hora, dependendo da concentração de pessoas. Não é possível prever se e quando vai acontecer uma crise numa multidão específica mas, supervisionando os comportamentos em grande escala, a transição para o movimento irregular pode ser detetada em diferentes partes de uma grande multidão e é possível tomar medidas para diminuir a concentração de pessoas nos pontos de tensão onde está a dar-se a transição que tende para o caos.

Sempre senti que um diamante oferecido brilha muito mais do que um comprado por nós.

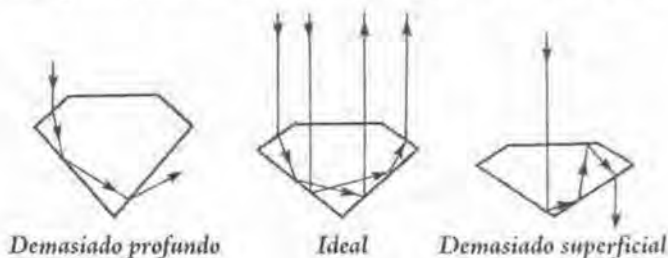
Mae West

Os diamantes são pedaços de carbono extraordinários. São o material mais duro que ocorre naturalmente no mundo.

As propriedades mais brilhantes dos diamantes são, contudo, óticas, porque o diamante tem um índice muito elevado de reflexão (2,4) em comparação com o da água (1,3) ou o do vidro (1,5). Isto significa que os raios de luz são curvados (ou "refratados") num ângulo muito grande quando atravessam um diamante. E o que é ainda mais importante, a luz que incide sobre a superfície do diamante a um ângulo superior a apenas 24 graus na vertical em relação à face deste será completamente refletida e não atravessará o diamante. Trata-se de um ângulo muito pequeno – para a luz que atravessa o ar ou a água este ângulo crítico é de cerca de 48 graus a partir da vertical, e para o vidro é de cerca de 42 graus.

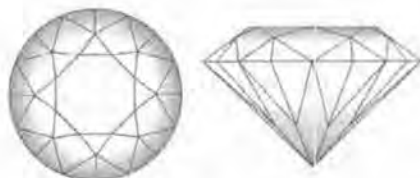
Os diamantes também espalham a cor de uma forma extrema. A luz branca comum é composta por um espectro de ondas de luz vermelha, laranja, amarela, verde, azul, índigo e violeta, que atravessam o diamante a velocidades diferentes e que sofrem curvaturas de ângulos diferentes (o vermelho curva-se à menor velocidade e o violeta à maior velocidade) enquanto a luz branca atravessa um meio transparente. O diamante produz uma diferença muito grande entre a maior e a menor curvatura das cores, chamada "dispersão", o que cria o célebre "fogo" multicolor quando a luz atravessa um diamante bem lapidado. Nenhuma outra pedra preciosa tem uma capacidade de dispersão tão grande. O desafio com que um joalheiro se depara é o de cortar o diamante de forma a que brilhe com tanta intensidade e cor quanto possível sob a luz que reflete para os olhos

do observador. A lapidação de diamantes é uma prática muito antiga, que já existe há milhares de anos, mas houve um homem que contribuiu mais do que qualquer outro para a nossa compreensão da melhor forma de cortar diamantes e do motivo pelo qual é a melhor. Marcel Tolkowsky (1899-1991) nasceu em Antuérpia no seio de uma família de lapidadores e comerciantes de diamantes. Era um jovem talentoso e depois de se ter licenciado na Bélgica, foi mandado para a Imperial College em Londres para estudar engenharia⁸. Enquanto ainda era apenas um estudante universitário, publicou um livro notável intitulado *Diamond Design*, que mostrou pela primeira vez como o estudo do reflexo e da refração da luz num diamante pode revelar a melhor forma de o cortar de modo a alcançar o máximo de brilho e "fogo". A elegante análise de Tolkowsky dos caminhos que são seguidos pelos raios de luz no interior de um diamante levou-o a propor um novo tipo de corte, o corte "Brilhante" ou "Ideal", que é atualmente o estilo usado preferencialmente para os diamantes redondos. Tolkowsky analisou os caminhos seguidos pelos raios de luz que incidiam em linha reta sobre a face do diamante e procurou descobrir quais os ângulos na parte inferior deste que deveriam ser inclinados de modo a refletir completa e internamente a luz no primeiro e segundo reflexos internos. Isto terá como resultado a refração quase total da luz que incide sobre o diamante, voltando a sair pela face deste e produzindo uma aparência mais brilhante. Para apresentar o máximo de brilho possível, os raios de luz que passam para o exterior não devem sofrer uma curvatura significativa da sua posição vertical quando se projetam para fora do diamante depois de refletidos internamente. As três imagens apresentadas em seguida mostram os efeitos de ângulos de corte demasiado grandes e demasiado pequenos, em comparação com o ângulo ideal que evita a perda de luz por dispersão através das faces inferiores e diminui o reflexo.

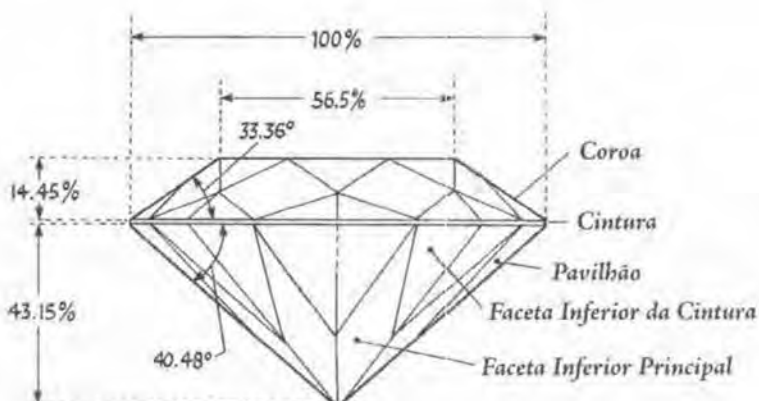


⁸ A sua tese de doutoramento foi sobre a moagem e polimento de diamantes e não sobre a aparência deles.

Tolkowsky estudou em seguida o equilíbrio ideal entre o brilho refletido e a dispersão do seu espectro de cores, de modo a criar um “fogo” especial e as melhores formas para as diferentes faces²⁰.



A sua análise, usando a matemática simples dos raios de luz, produziu uma receita para um diamante belo com um “corte brilhante”, com 58 facetas e um conjunto especial de proporções e ângulos com os valores ideais para produzir os efeitos visuais mais espetaculares quando o diamante é movido ligeiramente em frente aos nossos olhos. Mas iremos descobrir as que aparências enganam.



Neste diagrama vemos a forma clássica que Tolkowsky recomendou para um corte ideal, com os ângulos estreitos que otimizam o “fogo” e o brilho. As proporções das partes do diamante (indicadas com os seus nomes especiais) são apresentadas como percentagens do diâmetro da cintura, que é o diâmetro global*.

* A estreita espessura da cintura destina-se a evitar uma aresta muito vincada.

Porque Deus sabe que, no dia em que o comerdes, abrirem-se-ão os vossos olhos e sereis como Deus, ficando a conhecer o bem e o mal.

Livro do Gênesis

O ntem vi o filme *Eu, Robot*, baseado na obra do grande escritor de ficção científica Isaac Asimov. Em 1942, ele introduziu o conceito futurista em que os seres humanos coexistiam com robôs muito avançados num conto intitulado *Runaround*. Para garantir que os seres humanos não eram destruídos nem escravizados pelos seus assistentes perfeitamente eficientes, criou um conjunto de “Leis” que eram incluídas na programação dos cérebros eletrônicos de todos os robôs como medida de precaução. A escolha das leis que deveriam ser incluídas é muito interessante e não é apenas uma questão de saúde e segurança tecnológica, mas também uma questão mais profunda para todos aqueles que se perguntam por que existe Mal no mundo e que passos podem ser dados por uma divindade benevolente para o impedir.

As três leis originais de Asimov são baseadas nas três leis da termodinâmica.

Primeira Lei: Um robô não pode ferir um ser humano nem, por não agir, permitir que um ser humano seja ferido.

Segunda Lei: Um robô tem de obedecer às ordens que lhe são dadas pelos seres humanos, exceto nos casos em que essas leis entrem em conflito com a Primeira Lei.

Terceira Lei: Um robô tem de proteger a sua própria existência, desde que essa proteção não entre em conflito com a Primeira Lei nem com a Segunda Lei.

Mais tarde, Asimov acrescentou a "Lei Zero", mais uma vez baseando-se nas leis da termodinâmica, lei essa que se sobreporia à Primeira Lei:

Lei Zero: Um robô não pode prejudicar a Humanidade nem, por não agir, permitir que a Humanidade seja prejudicada.

O motivo da introdução desta última lei não é difícil de descobrir. Suponhamos que um louco obtinha acesso a um dispositivo nuclear que podia destruir o mundo inteiro e que só um robô podia impedi-lo de o acionar. Nesse caso, a Primeira Lei impedia o robô de agir para salvar a Humanidade. O problema da Primeira Lei é a falha em agir por parte dos robôs, mesmo quando a Lei Zero é irrelevante. Se o meu robô e eu ficassemos naufragados e fossemos parar a uma ilha deserta, e o meu pé com gangrena precisasse de ser amputado para me salvar a vida, o meu robô seria capaz de agir contra a Primeira Lei e amputá-lo? E um robô alguma vez poderia ser juiz num tribunal onde teria de atribuir penas àqueles que fossem considerados culpados por um júri?

Deveríamos sentir-nos seguros se fossem criados robôs em grandes quantidades, com estas quatro leis gravadas nos seus cérebros eletrônicos? Acho que não. É tudo uma questão de oportunidade. A precedência da Lei Zero sobre a Primeira significa que o robô pode matar-nos por estarmos a conduzir um carro com um consumo excessivo de combustível ou por não reciclarmos todas as embalagens que usamos. Ele considera que se este nosso comportamento continuar, pode pôr em perigo a Humanidade. Também pode levar muito a sério o seu dever de agir contra alguns líderes políticos do mundo. Pedir a um robô que tome medidas para preservar a Humanidade é perigoso porque lhe pede uma coisa que ainda não está definida. Não há uma única resposta para a pergunta: "O que é o Bem da Humanidade?" Não pode existir um computador que emita uma lista de todas as ações que são boas para a Humanidade e de todas as ações que lhe são prejudiciais. Nenhum programa informático pode ensinar-nos a distinguir o Bem do Mal.

É até provável que nos sintamos mais seguros sem a Lei Zero do que com ela. Ainda assim, existe outra consideração preocupante que pode pôr-nos em perigo de sofrer as ações diretas das quais a Primeira, a Segunda e a Terceira Leis procuram proteger-nos. Os robôs avançados terão

pensamentos complicados, pensamentos acerca de si próprios e de nós, tal como acerca de objetos inanimados: terão uma psicologia. Tal como acontece com os seres humanos, poderão ter dificuldade em entender esses pensamentos. E também poderão sofrer de alguns dos problemas psicológicos que vitimam os seres humanos. Tal como existem casos de seres humanos que sofrem delírios que os levam a pensar que são robôs, poderia dar-se o caso de um robô pensar que era um ser humano. Nesse caso, o robô poderia fazer o que lhe apetecesse, porque já não acreditava que as Quatro Leis da Robótica se aplicavam a si. Fortemente relacionado com este problema estaria o desenvolvimento de crenças religiosas ou místicas na mente dos robôs. Então e a Terceira Lei? Que existência robótica é que tem de ser preservada? A matéria do robô? A alma que ele julga que reside na máquina que o constitui? Ou a “ideia” do robô que continua a viver na mente do seu criador?

O leitor pode continuar a fazer perguntas deste tipo, mas perceberá que não é fácil evitar as consequências da inteligência artificial incluindo restrições e regras na programação. Quando aquela “coisa” a que chamamos “consciência” surge, as suas consequências são imprevisíveis e têm um imenso potencial tanto para o Bem como para o Mal, e é difícil ter um sem o outro – um pouco como na vida real.

79 | ROMPER COM AS CONVENÇÕES

Muitas pessoas prefeririam morrer a pensar; e, de facto, é o que a maioria delas faz.

Bertrand Russell

É fácil cairmos no erro de analisar um problema segundo os moldes estabelecidos. Sair desses moldes de pensamento e adotar uma abordagem “imaginativa” ou original à resolução do problema pode exigir uma forma diferente de pensar o problema e não a simples implementação correta dos princípios já aprendidos. Problemas simples que envolvem a aplicação de regras fixas de uma forma irrepreensível podem ser resolvidos usando a segunda abordagem. Por exemplo, se alguém o desafiar para um jogo do galo, existe uma forma de nunca perder, independentemente de ser o primeiro ou o segundo a começar. No entanto, esta estratégia, cujo resultado é um empate, só lhe garantirá a vitória se o seu adversário se desviar da estratégia ideal. Infelizmente, nem todos os problemas são tão fáceis como escolher a melhor jogada no jogo do galo. Eis um exemplo de um problema simples cuja solução irá quase de certeza deixá-lo surpreendido.

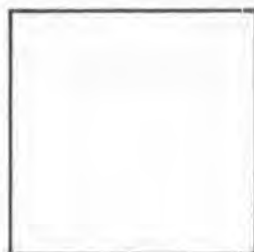
Desenhe um quadrado 3 x 3 composto por nove pontos. Em seguida, pegue num lápis e, *sem levantar o lápis do papel nem passar por cima de nenhuma linha já desenhada*, desenhe quatro linhas retas que unam todos os pontos.



Aqui está um exemplo de uma tentativa falhada. Não passa no ponto central no lado esquerdo.



Aqui está outra. Também deixa de fora um ponto, porque não atravessa o ponto central do quadrado.



Parece impossível, não parece? Consigo fazê-lo com quatro linhas retas, mas para isso tenho de repetir alguns dos traços, começando por uma linha diagonal e depois voltando atrás para desenhar três linhas cruzadas, só que isso obriga-me a fazer muito mais do que quatro linhas, apesar de o resultado final parecer mostrar apenas quatro:

80 | GOOGLE NAS CARAÍBAS: O PODER DA MATRIZ

O críquete é a prática de preguiça organizada.

William Temple

A maioria dos desportos cria tabelas de classificações para os campeonatos, que nos permitem apurar qual é a melhor equipa depois de todos os participantes se terem defrontado. A forma como a pontuação é atribuída às vitórias, derrotas e empates pode ser crucial para determinar quem fica no topo da tabela das classificações. Há alguns anos, as ligas de futebol decidiram atribuir três pontos às vitórias, em vez de dois, na esperança de que este sistema encorajasse um tipo de jogo mais ofensivo. Uma equipa obterá muito mais pontos por ganhar do que por empatar – situação em que cada uma das equipas recebe apenas um ponto. Mas este sistema parece bastante rudimentar. Afinal de contas, uma equipa não deveria receber muito mais pontos por derrotar uma equipa forte do que por derrotar uma das últimas classificadas?

A Taça Mundial de Críquete das Caraíbas em 2007 é um bom exemplo. Na segunda fase da competição, as oito equipas principais jogaram entre si (na verdade, cada uma delas jogou contra uma das outras que já estava na primeira fase e o resultado foi transportado para a frente, pelo que só tiveram de jogar mais seis jogos). Foram atribuídos dois pontos por cada vitória, um ponto por cada empate e zero por cada derrota. As quatro equipas no topo da tabela de classificações qualificaram-se para as duas eliminatórias das semifinais. No caso de várias equipas terem o mesmo número de pontos, eram separadas pela taxa de marcações global. Aqui fica a tabela:

Tabela de Classificações

EQUIPA	Jogos	Vitórias	Nulos	Derrotas	Pontos Individuais	Pontos
Austrália	7	7	0	0	2,40	14
Sri Lanka	7	5	0	2	1,48	10
Nova Zelândia	7	5	0	2	0,25	10
África do Sul	7	4	0	3	0,31	8
Inglaterra	7	3	0	4	-0,39	6
Índias Ocidentais	7	2	0	5	-0,57	4
Bangladesh	7	1	0	6	-1,51	2
Irlanda	7	1	0	6	-1,73	2

Mas consideremos outra forma de determinar as classificações das equipas que dê mais crédito às vitórias sobre uma equipa boa do que sobre uma equipa fraca. Damos a cada equipa uma pontuação que corresponde à soma dos pontos das equipas que derrotam. Uma vez que não houve empates, não temos de nos preocupar com eles. As pontuações globais são iguais a uma lista de oito equações:

$$A = SL + N + SA + E + IO + B + I$$

$$SL = N + IO + E + B + I$$

$$N = IO + E + B + I + SA$$

$$SA = IO + E + SL + I$$

$$E = IO + B + I$$

$$IO = B + I$$

$$B = SA$$

$$I = B$$

Esta lista pode ser expressa como equação matriz para a lista $\underline{x} = (A, N, IO, E, B, SL, I, SA)$ com a forma $A \underline{x} = K \underline{x}$, em que K é uma constante e A é uma matriz 8×8 de zeros e uns que denotam derrotas e vitórias, respetivamente:

	A	N	IO	E	B	SL	I	SA
A	0	1	1	1	1	1	1	1
N	0	0	1	1	1	0	1	1
IO	0	0	0	0	1	0	0	0
E	0	0	1	0	1	0	1	0
B	0	0	0	0	0	0	0	1
SL	0	1	1	1	1	0	1	0
I	0	0	1	0	1	0	0	0
SA	0	0	1	1	0	1	1	0

Para resolvermos as equações e encontrarmos a pontuação total de cada equipa, obtendo assim a sua classificação no campeonato segundo este sistema de pontuação diferente, temos de encontrar o autovector da matriz A com todas as entradas positivas ou zero. Cada uma destas soluções exigirá que K assuma um valor específico. Este corresponde a uma solução para a lista x na qual todos têm pontuações positivas (ou zero – se tiverem perdido todos os jogos), como é obviamente exigido pela situação aqui descrita. Ao resolvermos a matriz para este autovector, descobrimos que a solução é:

$$X = (A, N, IO, E, B, SL, I, SA) = (0,729; 0,375; 0,104; 0,151; 0,153; 0,394; 0,071; 0,332)$$

Neste caso, a classificação das equipas é determinada pelas magnitudes das suas pontuações, ficando a Austrália (A) no primeiro lugar com 0,729 e a Irlanda (I) em último com 0,071. Se compararmos estas classificações com a tabela original, temos:

Tabela de Classificações	Minha Tabela
A	A 0,729
SL	SL 0,394
N	N 0,375
SA	SA 0,332
E	B 0,153
IO	E 0,151
B	IO 0,104
I	I 0,071

As primeiras quatro equipes a qualificarem-se para as semi-finais ficam exatamente na mesma ordem em ambos os sistemas, mas três das últimas quatro obtêm uma classificação diferente. O Bangladesh ganhou apenas um jogo, obteve apenas dois pontos e ficou em penúltimo lugar na tabela de classificações oficial. No nosso sistema, passa para quinto lugar, porque a sua única vitória foi contra os Sul-africanos, que têm uma classificação muito superior. Inglaterra ganhou dois jogos, mas ambos contra as duas últimas classificadas e acabou a ficar abaixo do Bangladesh (ainda que separados apenas pela terceira casa decimal -- 0,153 vs. 0,151). As Índias Ocidentais ficaram em sexto lugar no sistema da Liga, mas desceram um lugar no nosso sistema.

O sistema de classificações aqui apresentado é o que está na base do motor de pesquisa Google. A matriz dos resultados quando a equipa *i* joga contra a equipa *j* corresponde ao número de ligações que existem entre o tópico *i* e o tópico *j*. Quando procuramos um termo, é criada uma matriz e "pontuações" pela gigantesca máquina informática que a Google tem ao seu dispor, e que resolve a equação da matriz para encontrar o autovector, criando assim a lista de correspondências para o termo que está a pesquisar. Ainda assim, continua a parecer magia.

81 | AVERSÃO À PERDA

Em teoria não há diferença entre a teoria e a prática. Na prática há.

Yogi Berra

As pessoas parecem reagir de formas muito diferentes à possibilidade de ganhar e perder. Os economistas demoraram muito tempo a reconhecer que o comportamento humano não é simétrico neste aspeto quando se trata de tomar decisões. Temos tendência para ter aversão ao risco e esforçamo-nos muito mais para evitar uma pequena perda do que para garantir um ganho significativo. Ter "aversão à perda" significa que perder uma nota de 50£ na rua produz uma quantidade de infelicidade superior à produzida por encontrar uma nota de 50£. Sentimo-nos melhor evitando uma taxa acrescida de 10% do que usufruindo de um desconto de 10% no preço de um bilhete de comboio.

Imaginemos que o leitor é um comerciante que vende produtos numa banca à beira da estrada. Decide que quer obter um determinado rendimento diário e que continuará a trabalhar até alcançar esse nível de vendas. O que acontece? Quando o dia lhe corre bem, atinge rapidamente o seu objetivo de vendas e vai para casa mais cedo. Quando o dia lhe corre mal, trabalha até cada vez mais tarde para atingir o seu objetivo. Isto parece irracional. Trabalha muito mais para evitar falhar a sua meta, mas não aproveita a oportunidade de trabalhar até mais tarde quando a procura é grande. Este é um exemplo clássico da psicologia da aversão ao risco.

Algumas pessoas argumentam que este tipo de comportamento é simplesmente irracional. Que não existe uma boa justificação para ele. Por outro lado, os ganhos e perdas não são necessariamente simétricos no que diz respeito à quantidade de dinheiro que possui atualmente. Se o valor total da sua riqueza for 100 000£, um ganho de 100 000£ é recebido

com muito agrado, mas uma perda de 100 000£ tem de ser evitada a todo o custo, porque o levará à bancarrota. A perda potencial é muito superior ao ganho potencial.

Por vezes, a tomada de decisões baseia-se numa perceção puramente psicológica de diferenças aparentes que, na realidade, não existem. Por exemplo, suponhamos que a Proteção Civil tinha de desenvolver planos para contrariar os efeitos nas casas costeiras de uma maré anormalmente alta e de uma tempestade que se previa que viesse a destruir 1000 casas. Os técnicos deste departamento pedem às pessoas que escolham um de dois planos. O plano A usa todos os recursos para construir um muro num determinado local, salvando 200 casas. O plano B usa os recursos de forma mais diversificada e salva todas as 1000 casas da destruição com uma probabilidade* de 1/5. Confrontadas com esta escolha, a maioria das pessoas escolhe o plano garantido e que inspira confiança, o A.

Agora, imaginemos que a Proteção Civil tem um Relações Públicas diferente que quer apresentar estes dois planos de outra forma. Neste caso, a escolha será entre o plano C, que permite que 800 casas sejam destruídas, e o plano D, que impede a destruição de todas as casas, com uma probabilidade de 1/5 e que permite a destruição de todas as casas com uma probabilidade** de 4/5. A maioria das pessoas escolhe o plano D. Isto é estranho, uma vez que o plano D é igual ao plano B e o plano A é igual ao plano C. A nossa aversão inata ao risco leva-nos a escolher o D em detrimento do C, mas não o B em detrimento do A, porque somos mais sensíveis às perdas. A perda garantida de 800 casas parece pior do que uma probabilidade de 4/5 de perder 1000. Mas no que diz respeito a salvar as casas, não respondemos tão fortemente à possibilidade de salvar 1000 casas como à certeza de salvar 200. Que coisa estranha.

* Isto significa que o número de casas que se espera que sejam salvas é $1000 \times 1/5 = 200$, o mesmo número salvo através do plano A.

** O número esperado de casas destruídas é 800 em ambos os planos, sendo o número de casas salvas 200, tal como nos planos A e B.

Somos todos lápis nas mãos de Deus.

Madre Teresa

O lápis moderno foi inventado em 1795 por Nicolas-Jacques Conté, um cientista que serviu no exército de Napoleão Bonaparte. O material mágico que era tão apropriado para a finalidade que pretendia servir era a forma de carbono puro a que chamamos grafite. Foi descoberta na Europa, na Baviera, no início do século xv, embora os Astecas já o usassem como marcador há várias centenas de anos. Inicialmente, acreditava-se que era um tipo de chumbo e chamavam-lhe *plumbago* ou chumbo negro. Só em 1789 é que começou a ser designado por grafite, baseada na palavra grega “graphein” que significa “escrever”.

Os depósitos mais puros de grafite foram encontrados em Borrowdale, perto de Keswick, no Lake District, em 1564, e originaram na zona uma grande atividade de contrabando e o mercado negro a ele associado. Durante o século xix, desenvolveu-se uma grande indústria de produção de lápis na zona de Keswick para explorar a enorme quantidade de grafite que ali existia. A primeira fábrica abriu em 1832 e a Cumberland Pencil Company acaba de celebrar o seu 175.º aniversário, apesar de as minas da zona terem fechado há muito tempo e de a grafite usada vir agora do Sri Lanka e de outros locais distantes. Os lápis da Cúmbria eram os que tinham mais qualidade porque a grafite usada deitava pó e marcava muito bem o papel. O processo original de Conté para a produção de lápis envolvia cozer uma mistura de água, barro e grafite num forno a 1040° C antes de envolver em madeira o material sólido obtido. A forma desse invólucro podia ser quadrada, poligonal ou redonda, dependendo da utilização a que se destinava o lápis – por exemplo, um carpinteiro não tem interesse

num lápis que possa rolar para fora da bancada de trabalho. A dureza da mina final do lápis pode ser determinada ajustando as frações relativas de barro e grafite na mistura que é cozida. Os fabricantes vendem habitualmente lápis com 20 níveis distintos de dureza, do mais macio, 9B, ao mais duro, 9H, sendo o valor intermédio mais popular, o HB, algures entre o H e o B. H significa "hard", duro, e B significa "black", negro. Quanto maior o número dentro da categoria B, mais grafite fica depositada no papel. Também existe a categoria F, ou ponta fina, que é um lápis duro, adequado para a escrita, mas não para o desenho.

A característica mais estranha da grafite é ser uma forma de carbono puro que é um dos sólidos mais moles conhecidos e também um dos melhores lubrificantes, por causa dos seus átomos de carbono que se unem para formar um aro que pode deslizar sobre aros adjacentes. No entanto, se mudarmos a estrutura atômica, encontramos uma outra forma cristalina do carbono puro, o diamante, que é um dos sólidos mais duros que se conhecem.

Será interessante perguntar qual o comprimento de uma linha que se pode desenhar com um lápis HB comum até se gastar completamente a mina. A espessura da grafite depositada numa folha de papel por um lápis macio 2B é de cerca de 20 nanómetros e um átomo de carbono tem um diâmetro de 0,14 nanómetros, pelo que a linha do lápis tem apenas 143 átomos de espessura. A mina do lápis tem um raio de cerca de 1 mm, tendo assim uma área de π mm quadrados. Se o comprimento do lápis for 15 cm, o volume de grafite que será depositado numa linha reta é de 150π mm cúbicos. Se desenharmos uma linha com 20 nanómetros de espessura e 2 mm de largura, a mina será suficiente para continuar ao longo de uma distância de $L = 150\pi/4 \times 10^{-7}$ mm = 1,178 quilómetros. No entanto, não teste esta minha previsão!

83 | TESTE DE RESISTÊNCIA DO ESPARGUETE

Sempre que vir uma carrinha da Parceline, vou lembrar-me do Miles Klingron. Porque foi o Miles quem decidiu que aquele nome era adequado para um prato de comida italiana.

Richard Ingram

Segure em ambas as pontas de um fio de esparguete, longo, seco e quebradiço. Curve-o e aproxime gradualmente as duas pontas até o partir. Provavelmente esperava que ele se partisse em dois pedaços, deixando-o com um em cada mão. Estranhamente, isso nunca acontece. O esparguete parte-se sempre em mais do que dois pedaços. Se partir um pedaço fino de madeira ou de plástico, este divide-se em dois pedaços. Por que é que o comportamento do esparguete é diferente? Richard Feynman também ficou intrigado com esta questão e, na sua biografia, que Daniel Hills escreveu, conta-se a seguinte história:

Um dia, estávamos a fazer esparguete... Se pegarmos num fio de esparguete e o partirmos, constataremos que se dividirá quase sempre em três pedaços. A que se deve isto? Por que é que se parte em três pedaços? Passámos quase duas horas a inventar teorias loucas. Inventámos experiências, como quebrá-lo debaixo de água, pensando que esta podia abafar o som e consequentemente as vibrações. Ao fim de duas horas, acabámos com esparguete partido espalhado por toda a cozinha e sem uma boa teoria que explicasse por que motivo o esparguete se parte sempre em três pedaços.

Recentemente, o problema, que acabou por se revelar inesperadamente difícil, foi resolvido. Um pedaço quebradiço de qualquer material, não apenas de esparguete, partir-se-á quando for fletido para lá da sua curvatura crítica, chamada a "curvatura de rutura". Não há mistério algum nisto, mas o que acontece em seguida é bastante interessante. Quando se

dá a quebra, uma extremidade de cada parte fica livre ao passo que a outra ainda está na sua mão. A extremidade livre, que se viu subitamente libertada, tenta endireitar-se e envia ondas de curvatura ao longo do seu comprimento, na direção da mão, onde está presa. As ondas são refletidas e encontram-se com outras que chegam a diferentes pontos do fio de esparguete. Quando se encontram, dá-se um salto súbito na curvatura, o suficiente para voltar a partir o fio de esparguete fletido. Novas ondas de curvatura são produzidas por esta nova quebra e podem conduzir a novos aumentos localizados da curvatura para além do valor crítico, em pontos diferentes do esparguete. Consequentemente, o esparguete irá quebrar-se num ou mais pontos depois da primeira quebra. As quebras param quando já não restar energia suficiente para permitir que as ondas se desloquem ao longo do fio de esparguete que ficou na sua mão. Todos os fragmentos que ficam livres em ambas as pontas caem para o chão.

Ideias novas, frescas como pepínos.

Stephen Moss

A construção moderna mais dramática na cidade de Londres é o edifício do número 30 da rua St. Mary Axe, mais conhecido como o edifício da seguradora suíça Swiss Re, a pinha ou, simplesmente, o Gherkin*. O Príncipe Carlos vê-o como sintomático de uma praga de torres que mancham o rosto de Londres. Os arquitetos que o desenharam, Norman Foster and Partners, apresentaram-no como edifício emblemático dos tempos modernos e receberam o prémio RIBA Stirling em 2004 pela sua criação. O edifício alcançou o objetivo de dar visibilidade à seguradora Swiss Re e estimulou um grande debate acerca da aceitabilidade de torres nos horizontes tradicionais da cidade de Londres. Infelizmente, embora continue a haver um intenso debate acerca do sucesso estético do Gherkin, não há muitas dúvidas de que foi uma desilusão para a Swiss Re, do ponto de vista comercial. A empresa ocupa apenas os primeiros 15 de 34 pisos, e nunca conseguiu alugar a outra metade do edifício a nenhuma outra empresa. Este facto não é totalmente surpreendente: o tipo de empresa de grandes dimensões que pode pagar o aluguer de um espaço daquele tipo considera que o edifício ficou de tal forma associado à Swiss Re que a empresa arrendatária permaneceria para sempre em segundo plano e não tiraria benefício algum da sua presença ali. Assim, o espaço foi arrendado em frações menores.

O traço mais óbvio do Gherkin é a sua altura – 180 metros – e a criação de uma torre com uma escala tal que cria problemas estruturais e ambientais. Hoje em dia, os engenheiros podem criar sofisticados modelos

* Vegetal semelhante a um pepíno. (N. T.)

informáticos de um grande edifício que lhes permitem estudar a reação deste ao vento e ao calor, a quantidade de ar fresco que recebe do exterior e o seu efeito nos transeuntes ao nível da rua. A alteração de um dos aspetos do *design*, como o nível de refletividade da sua superfície, afetará muitas outras áreas – como, por exemplo, a temperatura interna e as necessidades relativas ao ar condicionado – e todas essas consequências podem ser observadas imediatamente, usando sofisticadas simulações informáticas do edifício. De nada serve seguir a abordagem “uma coisa de cada vez” na conceção de uma estrutura complicada como um edifício moderno – é preciso fazer várias coisas ao mesmo tempo.

O elegante perfil curvo do Gherkin não tem apenas uma função estética nem é simplesmente reflexo do desejo de um louco de criar uma obra espetacular e controversa. A forma cónica, que começa por ser mais estreita ao nível do chão e alcança a máxima largura no piso 16, antes de voltar a estreitar a um nível regular em direção ao topo, foi escolhida em resposta às simulações por computador.

Os edifícios altos canalizam os ventos à sua volta ao nível da rua (de uma forma semelhante a tapar parcialmente com o dedo a abertura de uma mangueira para fazer a água espalhar-se por uma área maior – o aumento da pressão produzido por esta compressão resulta numa maior velocidade do fluxo de água) e isto pode ter um efeito terrível nos transeuntes e nas pessoas que estão a usar o edifício. A sensação é semelhante à de estar no interior de um túnel de vento. O estreitamento do edifício na base reduz estes efeitos indesejados do vento, por haver menos compressão das correntes de ar. O estreitamento no topo também tem um papel importante. Se se colocar ao lado de uma torre convencional e não cónica ao nível do chão e olhar para cima, o edifício parecer-lhe-á gigante e bloqueará a visibilidade para uma grande parte do céu. Uma configuração cónica deixa a descoberto uma maior parte do céu e reduz o efeito dominador da estrutura, por não se conseguir ver o topo ao nível do chão.

Outra característica notável do exterior do edifício é o facto de ser redondo e não quadrado ou retangular. Mais uma vez, é uma característica vantajosa para suavizar e abrandar o fluxo do ar em torno do edifício. Também ajuda a tornar o edifício invulgarmente ecológico. Seis enormes triângulos foram recortados em cada piso, de fora para dentro. Permitem a entrada da luz e ventilação natural no interior do edifício, reduzindo a

necessidade de utilização de ar condicionado convencional e duplicando a eficiência energética do edifício em comparação com um edifício típico à mesma escala. Estes recortes não estão colocados imediatamente abaixo uns dos outros entre pisos, estando antes ligeiramente desalinhados dos que estão nos pisos acima e abaixo deles. Esta característica ajuda a aumentar a eficiência da sucção do ar para o interior. É este ligeiro desalinhamento dos seis recortes entre cada piso que cria o característico padrão em espiral no exterior.

Observando à distância o exterior arredondado do edifício, poderemos pensar que os painéis individuais da superfície são curvos – o que tornaria a construção muito cara –, mas na realidade não o são. Os painéis são suficientemente pequenos, em relação à distância através da qual a curvatura se torna significativa, para um mosaico de painéis lisos com quatro lados ser suficiente para desempenhar a sua função. Quanto mais pequenos forem, melhor conseguirão aproximar a cobertura da superfície curva exterior. Todas as mudanças de direção são feitas nos ângulos que unem os painéis,

85 | CALCULAR O ÍNDICE DE PREÇOS

Em média, um ser humano tem um seio e um testículo.

Des McHale

Todos os países economicamente desenvolvidos têm algum nível de alteração do custo de vida do cidadão comum que deriva do preço das unidades padrão de um conjunto de bens representativos, incluindo alimentos básicos, leite, aquecimento e luz. Têm nomes como Índice dos Preços no Consumidor (IPC) e são um dos indicadores tradicionais da inflação que, por sua vez, pode ser usado para indexar os salários e os subsídios. Assim, para os cidadãos é desejável que este indicador seja alto, ao passo que aos governos interessa que seja baixo.

Uma das formas de determinar um índice de preços é simplesmente calcular a média de um conjunto de preços – basta somá-los e dividi-los pelo número de preços somados. Isto é o que os estatísticos chamam a *média aritmética*, ou simplesmente uma “média”. Normalmente, queremos descobrir como as coisas estão a mudar ao longo do tempo – o custo do mesmo conjunto de bens está a subir ou a descer de mês para mês? – pelo que é preciso comparar a média do ano anterior com o seu valor neste ano, dividindo o primeiro pelo segundo. Se o resultado for superior a 1, os preços estão a descer; se for inferior a 1, estão a aumentar. Isto é bastante simples, mas haverá algum problema oculto?

Suponhamos que uma família gasta habitualmente o mesmo dinheiro todas as semanas em carne de vaca e peixe, e que depois o preço da carne de vaca duplica ao passo que o preço do peixe se mantém. Se continuarem a comprar a mesma quantidade de carne de vaca e de peixe, o valor total que gastam em carne e peixe será 1,5 vezes superior ao que gastavam

anteriormente – um aumento de 50%. A média das subidas de preços será $\frac{1}{2} \times (1+2)=1,5$. O $\frac{1}{2}$ está apenas a dividir o número de produtos (dois: carne de vaca e peixe); 1 é o fator pelo qual o preço do peixe muda (mantém-se) e 2 é o fator pelo qual o preço da carne de vaca muda (duplica).

Este fator de inflação de 1,5, um aumento de 50%, será a estatística principal. Mas para uma família que não coma carne e que, consequentemente, não compre carne de vaca, não terá qualquer importância. Se só consumissem peixe, as suas despesas semanais não sofreriam qualquer alteração. O fator da inflação também é uma média de todas as escolhas alimentares das famílias. Baseia-se numa suposição relativa à psicologia humana: o pressuposto de que a família continuará a comer a mesma quantidade de carne e de peixe, apesar do aumento do custo relativo da carne em comparação com o do peixe. Na realidade, as famílias podem ter comportamentos diferentes e decidir ajustar a quantidade de carne e peixe que comprem, de modo a gastarem o mesmo dinheiro em cada um dos produtos. Isto significaria que no futuro comprariam menos carne por causa do seu aumento de preço.

A suposição de que as famílias reagem às mudanças de preço mantendo constante a fração do seu orçamento que gastam em cada produto sugere que a média aritmética simples usada no cálculo do índice de preços deveria ser substituída por outro meio.

A *média geométrica* de duas quantidades é a raiz quadrada do seu produto*. Desta forma, o índice obtido pelo cálculo da média geométrica para a mudança do preço da carne sem alteração ao preço do peixe é:

$$\begin{aligned} & \sqrt{(\text{Novo preço da carne/preço antigo da carne}) \times (\text{Novo preço do peixe/preço antigo do peixe})} \\ & = \sqrt{1 \times 2} = 1,41 \end{aligned}$$

O que estas duas formas de medir a inflação têm de interessante é que a média geométrica de qualquer número de quantidades nunca é superior à média aritmética²¹, pelo que os governos preferem, obviamente, a média geométrica²². Esta sugere uma inflação mais baixa e resulta num índice de inflação mais baixo para os aumentos dos salários e pensões.

* Em termos mais gerais, a média geométrica de n quantidades é a n ésima raiz do seu produto.

²² Nos Estados Unidos, o Ministério do Trabalho mudou a fórmula de cálculo do Índice dos Preços no Consumidor da média aritmética para a média geométrica em 1999.

Outro benefício da média geométrica é de natureza prática e não propriamente política. Para determinar a taxa de inflação, compara-se o índice em diferentes alturas. Se estiver a usar a média aritmética, tem de calcular o seu valor em 2008 e em 2007 para descobrir em que medida o “cabaz” de preços aumentou de preço no ano passado. No entanto, a média aritmética envolve a *soma* de todo o tipo de coisas que podem ser medidas em diferentes unidades: € por kg, € por litro, etc.; alguns destes termos envolvem o preço por unidade de peso, outros o preço por unidade de volume. É uma mistura, o que é um problema porque para calcular a média aritmética é preciso somá-los, o que por sua vez não faz sentido se as unidades em que são medidos diferentes tipos de bens forem diferentes. A média geométrica, por outro lado, permite-nos usar as unidades que quisermos para os preços dos bens em 2008, desde que sejam usadas as mesmas para os mesmos bens em 2007. Quando se divide o valor da média geométrica de 2008 pelo seu valor no ano de 2007, para encontrar o fator de inflação, todas as unidades diferentes se cancelam mutuamente por serem exatamente as mesmas tanto no denominador como no numerador. Como vê, é um índice muito eficaz.

86 | A OMNISCÊNCIA PODE SER UMA DESVANTAGEM

À venda pela proprietária, *Encyclopaedia Britannica* em ótimo estado. Já não é necessária. O marido sabe tudo.

Pequeno anúncio

Imagine como será saber tudo. Talvez não seja assim tão fácil. Talvez seja mais fácil imaginar como seria saber tudo o que gostaria de saber – ou que precisa de saber. E até isso parece excessivo: saber os números premiados da lotaria da próxima semana, o comboio que deve apanhar para evitar atrasos, quem vai ganhar aquele jogo de futebol importante. Isto tem grandes vantagens, embora a vida possa vir a revelar-se muito infeliz sem o benefício das ocasionais surpresas agradáveis.

Há um estranho paradoxo acerca da omnisciência que mostra que uma pessoa nessa situação pode estar pior do que uma pessoa que não saiba tudo. Suponhamos que está a ver um desafio entre dois pilotos de aviões que, à semelhança das justas, avançam na direção um do outro a grande velocidade. Perde o piloto que primeiro se desviar para o lado. Qual deverá ser a tática de um jogador num desafio deste tipo? Se nunca se desviar, pode acabar por morrer se o outro jogador tiver adotado a mesma tática – e nesse caso ninguém ganha. Se se desviar sempre, nunca ganha – podendo apenas empatar quando o outro piloto também se desviar. Claramente, desviar-se sempre é a única estratégia garantida para minimizar as perdas. Uma estratégia mista, desviando-se algumas vezes e escolhendo não se desviar noutras, produzirá algumas vitórias, mas acabará por resultar na morte do piloto, a não ser que o segundo piloto se desvie sempre, mesmo quando o primeiro não se desvia. Se o outro piloto pensar da mesma forma (ele analisa as suas possibilidades independentemente e não conhece a decisão do seu adversário), deverá chegar às mesmas conclusões.

Agora imagine-se a jogar este jogo contra um adversário omnisciente. Ele sabe qual vai ser a sua estratégia, portanto o leitor deveria escolher *nunca* se desviar. Ele saberia que a sua estratégia era *nunca* se desviar e escolheria desviar-se sempre. Desta forma, o piloto omnisciente nunca poderia ganhar!

Esta história pode ser aplicada ao mundo da espionagem. Se estiver a ouvir todas as comunicações dos seus inimigos e eles souberem que o leitor está a ouvir, a sua omnisciência pode deixá-lo numa posição de desvantagem.

87 | POR QUE É QUE AS PESSOAS NÃO SÃO MAIS ESPERTAS

Que Deus me ajude na minha busca da verdade e que me proteja daqueles que pensam tê-la encontrado.

Antiga oração inglesa

Quando os astrónomos especulam acerca da natureza de extraterrestres muito evoluídos ou quando os biólogos contemplam um futuro em que os seres humanos evoluíram até serem mais inteligentes do que atualmente, supõe-se sempre que este aumento da inteligência é uma coisa boa. O processo evolucionário aumenta a probabilidade de transmitir às gerações seguintes os traços que aumentam as nossas hipóteses de sobrevivência e de reprodução. É-nos difícil imaginar o risco que poderia derivar de um aumento exacerbado da inteligência da nossa espécie.

Se alguma vez teve a experiência de tentar gerir uma comunidade de indivíduos mais inteligentes do que a média, poderá ser da opinião contrária. Um bom exemplo de uma situação deste tipo poderá ser o diretor de um departamento de uma universidade ou o editor de um livro que está a ser escrito por vários autores. Em situações assim percebemos rapidamente que este tipo de inteligência elevada pode ser acompanhado de uma tendência para ser individualista, para pensar de forma independente e para discordar de todas as pessoas que têm opiniões diferentes. Talvez a capacidade de se relacionar com os outros, de colaborar com os outros em vez de trabalhar contra eles, tenha sido mais importante durante as primeiras fases do desenvolvimento da inteligência. Se a inteligência evoluir agora para níveis sobre-humanos, não haverá a possibilidade de os efeitos dessa evolução serem desastrosos do ponto de vista social? Por outro lado, níveis inferiores de inteligência são claramente uma desvantagem no que respeita a lidar com perigos potencialmente previsíveis. Poderá haver um nível ideal de inteligência para a vida num determinado ambiente que maximize as nossas possibilidades de sobrevivência a longo prazo.

A arte "transporta-nos"; o *design* não, a menos que seja o *design* de um autocarro.

David Hockney

Uma vez vi dois turistas a tentarem orientar-se nas ruas do centro de Londres usando como referência o mapa do metro. Embora seja uma opção ligeiramente melhor do que usar como referência o tabuleiro do Monopólio, não é especialmente útil. O mapa do metro de Londres é uma maravilhosa peça de *design* funcional e artístico, mas tem uma peculiaridade: as estações não estão colocadas no local geográfico correto. É um mapa *topológico*: apresenta de forma exata as ligações entre as estações, mas distorce a sua localização exata por motivos estéticos e práticos.

Quando Harry Beck propôs este tipo de mapa para a gestão do metro-politano de Londres, era apenas um jovem desenhador com conhecimentos de eletrónica. O Underground Railway foi formado em 1906, mas por volta da década de 1920 estava a ser um fracasso comercial, especialmente por causa da aparente extensão e complexidade das deslocações dos pontos mais distantes até ao centro de Londres, em especial nos casos em que era necessário trocar de linha. Um mapa geograficamente exato parecia extremamente confuso por causa da natureza desordenada das ruas do centro da cidade que se desenvolveram ao longo de centenas de anos sem qualquer tipo de planeamento central. Não era como Nova Iorque nem como Paris, com um plano organizado das ruas. As pessoas ficaram desmotivadas.

O elegante mapa de Beck, criado em 1931 – tendo sido inicialmente recusado pelo departamento de publicidade do metropolitano e pelo seu diretor-executivo, Frank Pick –, resolveu muitos problemas de uma só vez. Ao contrário do anterior mapa dos transportes públicos e à semelhança de um circuito eletrónico, usava apenas linhas horizontais, verticais

e diagonais em ângulos de 45 graus. Acabou por também incluir uma representação simbólica do rio Tamisa e uma forma agradável de representar as trocas de estação e a geografia distorcida dos arredores de Londres, fazendo lugares distantes como Rickmansworth, Morden, Uxbridge e Cockfosters parecerem próximos do centro da cidade. Beck continuou a aperfeiçoar e a aumentar o seu mapa durante os 40 anos seguintes, incluindo novas linhas e extensões de linhas antigas, procurando sempre torná-lo simples e claro. O mapa era sempre referido como o Diagrama do Metropolitano de Londres, ou simplesmente "O Diagrama", para evitar possíveis confusões com mapas tradicionais.

O *design* clássico de Beck tornou-se o primeiro mapa topológico. Pode ser modificado esticando-o ou distorcendo-o de qualquer forma que não quebre as ligações entre as estações. Imagine o mapa desenhado numa folha de borracha, que pode ser esticada ou torcida em todas as direções sem se partir ou rasgar. O impacto deste mapa foi sociológico, para além de cartográfico. Redefiniu a forma como as pessoas viam a cidade de Londres. Incluía no mapa os lugares mais afastados, fazendo os seus habitantes sentirem que estavam perto do centro de Londres. Até definiu os níveis dos preços das casas. Para a maioria de nós, é assim que Londres "é".

A ideia original de Beck faz muito sentido. Quando estamos no metro, debaixo do chão, não precisamos de saber onde nos encontramos, como acontece nas deslocações a pé ou de autocarro. A única coisa que importa é onde saímos e onde mudamos para outras linhas. Aproximar os lugares distantes do centro não só contribuiu para que os habitantes de Londres se sentissem menos isolados, como também ajudou a criar um diagrama com um aspeto equilibrado, que cabe perfeitamente numa folha dobrada que podemos guardar no bolso do casaco.

89 | NÃO HÁ NÚMEROS DESINTERESSANTES

Tudo é belo à sua maneira.

Ray Stevens

A lista dos números é infinita. Os números pequenos, como 1, 2 e 3, são usados constantemente para descrever os pequenos números da nossa vida: o nosso número de filhos, de carros, de itens numa lista de compras. O facto de haver tantos nomes para designar grupos de pequenas quantidades, que estão intimamente relacionados com as suas identidades – como, por exemplo, dobro, gémeos, braça, par, duo, casal, dueto – sugere que as origens destes são anteriores ao nosso sistema de contagem decimal. Cada um destes números pequenos parece ser interessante de alguma forma. O número 1 é o mais pequeno de todos, o 2 é o primeiro número par, o 3 é a soma dos dois números anteriores, o quatro é o primeiro número que não é primo e que pode, desta forma, ser dividido por outro número que não ele mesmo, 5 é a soma de um quadrado (2^2) mais 1. E a lista continua. Gradualmente, começamos a perguntar-nos se haverá algum que seja completamente desinteressante, que passe despercebido no meio de todos os outros como um número feio e apagado no grande baile dos números.

Podemos provar que esse número existe? Sim, podemos, se abordarmos esta questão da mesma forma que abordamos muitas outras questões matemáticas. Começamos por supor que o inverso é verdade e depois usamos a suposição para deduzir um facto que o contradiga. Isto significa que a suposição inicial era provavelmente falsa. É uma versão extrema da manobra de xadrez em que um jogador oferece uma peça ao adversário sabendo que, se este a apanhar, abrirá o caminho para um ganho muito

maior no futuro. É a versão extrema desta manobra porque é todo o jogo, e não apenas uma única peça, que está a ser oferecido nesta manobra lógica.

Suponhamos que existem números inteiros positivos desinteressantes e que vamos reuni-los a todos. Se um conjunto desse tipo puder existir, terá de ter um número mais pequeno. Mas esse número mais pequeno é, por definição, interessante: é o número desinteressante mais pequeno. Isto contradiz a nossa suposição inicial de que é um número desinteressante. Desta forma, a nossa suposição de que existem números desinteressantes era falsa. Todos os números devem, portanto, ser interessantes.

Eis uma história, muito conhecida entre matemáticos, que prova o que digo. Fala do matemático inglês Godfrey Hardy, quando foi visitar o seu amigo, o notável matemático indiano Srinivasa Ramanujan, num hospital em Londres: "Enquanto estava no táxi, reparou no número de identificação deste: 1729. Deve ter ficado a pensar naquele número, porque entrou no quarto de hospital onde estava Ramanujan, mal o cumprimentando, e começando imediatamente a expressar a sua desilusão com aquele número. Era, segundo declarava, 'um número muito aborrecido', acrescentando que esperava que não fosse um mau presságio. 'Não, Hardy', disse Ramanujan, 'é um número muito interessante. É o número mais pequeno que pode ser expresso como a soma de dois cubos de duas formas diferentes'". Esses números são atualmente conhecidos como números "de táxi" em memória deste episódio.

* Porque $1729 = 13 + 123 = 93 + 103$. Nestes exemplos os cubos têm de ser positivos. Se permitirmos que números negativos sejam elevados ao cubo, o número mais pequeno será $91 = 63 + (-53) = 43 + 33$.

O número que marcou não existe.

Mensagem gravada para números de telefone não atribuídos em Itália.^{*}

Nos séculos XVI e XVII não era invulgar os matemáticos mais famosos da época publicarem as suas descobertas em código. Isto parece muito estranho aos olhos dos cientistas modernos, que lutam pelo reconhecimento e por serem os primeiros a descobrir alguma coisa, mas havia método nesta aparente loucura dos matemáticos desses tempos remotos. Como se costuma dizer, queriam ter o bolo e comê-lo. Publicar uma descoberta que faz uso de um novo “truque” matemático estabelece esse cientista como o descobridor do truque, mas também o revela, permitindo que outros o usem e que lhe passem à frente em futuras, e possivelmente maiores, descobertas. O cientista depara-se com uma escolha: não anuncia imediatamente a descoberta até ter tido tempo de investigar as outras possibilidades mais pormenorizadamente – correndo o risco de outra pessoa descobrir e publicar a sua primeira descoberta – ou pode publicar uma versão codificada. Supondo que ninguém consegue desvendar o código, o novo truque está protegido da exploração por parte dos outros, mas se aparecer outro cientista a anunciar que descobriu o que ele já descobriu, pode simplesmente decodificar a mensagem para provar que fez aquela descoberta muito tempo antes. É um truque muito astuto – devo acrescentar que esse tipo de comportamentos não é bem visto nos campos da ciência e da matemática hoje em dia, e provavelmente não seria tolerado se alguém o tentasse fazer. No entanto, é algo que acontece no mundo da literatura. Livros como o romance político *Cores Primárias*, acerca da campanha de Bill Clinton para as eleições presidenciais, escrito por um

^{*} Em italiano, “Il numero selezionato da lei è inesistente”.

jornalista inicialmente não identificado a escrever sob pseudônimo, parece um pouco uma tentativa de ter o bolo e comê-lo.

Suponhamos que o leitor queria proteger a sua identidade de forma semelhante nos tempos de hoje – de que forma pode usar um raciocínio matemático simples para o conseguir? Basta escolher dois números primos muito grandes, por exemplo 104 729 e 105 037 (na verdade deve escolher números muito maiores, com centenas de dígitos, mas estes são suficientemente grandes para lhe dar uma ideia geral). Em seguida multiplique-os para obter o seu produto: 11 000 419 973. Não confie na calculadora, pois ela provavelmente não consegue processar um número tão grande e acabará por arredondar o resultado – a minha calculadora deu um resultado errado: 11 000 419 970.

Agora voltemos ao nosso desafio da publicação anónima. O leitor quer publicar a sua descoberta, mas não quer revelar publicamente a sua identidade, querendo, no entanto, incluir uma “assinatura” oculta, para no futuro poder provar que é o autor. Pode publicar o livro com esse grande produto dos dois números primos impresso na contracapa. Se conhecer os dois fatores (104 729 e 105 037) pode facilmente multiplicá-los e mostrar que a resposta é o seu “código”. No entanto, se outra pessoa tentar partir do resultado – 11 000 419 973 –, não terá tanta facilidade em encontrar os dois fatores. Se tivesse escolhido multiplicar dois números primos muito grandes, cada um com 400 dígitos, outra pessoa poderia demorar a vida toda a tentar encontrar os dois fatores, mesmo com a ajuda de um computador muito rápido. Não é impossível desvendar o nosso “código” mas se usarmos números muito grandes, demora muito tempo.

Esta operação de multiplicação e fatorização de números é um exemplo de uma operação “alçapão” (ver Capítulo 27). É rápido e fácil seguir numa direção (como cair por um alçapão), mas é demorado e difícil fazer o caminho inverso (como tentar voltar a subir por um alçapão). Uma versão mais complicada da multiplicação de dois números primos é usada como base da maioria dos atuais códigos comerciais e militares. Por exemplo, quando faz compras pela Internet e introduz o seu número de cartão de crédito num *website* seguro, os dados são combinados com grandes números primos, que são transmitidos à empresa e depois decodificados através da fatorização de números primos.

91 | O PARADOXO DA PATINAGEM NO GELO

Depois de acabar de jantar, Sidney Morgenbesser decide pedir a sobremesa. A empregada de mesa diz-lhe que tem duas opções: tarte de maçã e tarte de amora. Sidney pede uma tarte de maçã. Ao fim de alguns minutos, a empregada volta e diz-lhe que também têm tarte de cereja, ao que Morgenbesser responde: "Nesse caso, quero a tarte de amora!"

Lenda académica

Quando fazemos escolhas ou votos, parece racional esperar que se escolhamos K como a melhor entre todas as alternativas disponíveis e depois alguém vier oferecer-nos uma terceira alternativa, a opção Z, que se tinham esquecido de incluir na oferta inicial, a nossa nova escolha seja manter a opção K ou optar antes por Z. Qualquer outra escolha parece irracional, porque estaríamos a escolher uma das opções que rejeitámos inicialmente em favor da opção K. Como é que a oferta de uma nova opção pode mudar a ordem de preferências das opções anteriormente disponíveis?

A noção de que isto não devia poder acontecer está tão enraizada nas mentes da maioria dos economistas e matemáticos que é geralmente excluída por princípio na conceção dos sistemas de voto. No entanto, sabemos que a psicologia humana poucas vezes é completamente racional e que há situações em que a alternativa irrelevante muda a ordem das nossas preferências, como foi o caso da sobremesa de Sidney Morgenbesser (embora entre pedidos ele possa ter visto uma das tartes em questão).

Um exemplo célebre era um sistema de transportes que oferecia um serviço de autocarros vermelhos como alternativa ao carro. Aproximadamente metade das pessoas começou a usar o autocarro vermelho; a outra metade continuou a usar o carro. A seguir, foi apresentado um segundo serviço de autocarros, desta vez azuis. Seria de esperar que um quarto

das pessoas usasse o autocarro vermelho, que outro quarto usasse o autocarro azul e que metade continuasse a usar o carro. Por que haveriam de dar importância à cor do autocarro? Na realidade, o que aconteceu foi que um terço das pessoas começou a usar o autocarro azul, outro terço o autocarro vermelho e um terço a usar o carro!

Há uma famosa situação em que o efeito de alternativas irrelevantes foi realmente incluído no processo de seleção, com resultados tão bizarros que levaram os seus autores a abandonarem completamente esse processo. A situação em questão era a avaliação de provas de patinagem no gelo nos Jogos Olímpicos de inverno de 2002, em que a jovem americana Sarah Hughes derrotou as favoritas Michelle Kwan e Irina Slutskaya. Quando vemos provas de patinagem artística na televisão, as pontuações de cada prova individual (6.0, 5.9, etc.) são anunciadas com grande entusiasmo. No entanto, curiosamente, não são essas pontuações que determinam quem ganha. São apenas usadas para ordenar as patinadoras. O leitor poderá ter imaginado que o júri se limitava a somar todas as pontuações dos dois programas (curto e longo) executados por cada patinadora individual e que a que tinha a maior pontuação ganhava a medalha de ouro. Infelizmente, não foi o que aconteceu em 2002 em Salt Lake City. No fim do programa curto, a ordem das classificações das primeiras quatro concorrentes era:

Kwan (0.5), Slutskaya (1.0), Cohen (1.5), Hughes (2.0)

São-lhes atribuídas automaticamente as pontuações 0.5, 1.0, 1.5 e 2.0 por terem ficado nos primeiros quatro lugares da tabela de classificações (sendo a pontuação mais baixa a melhor). É de notar que todos os fabulosos 6.0 foram esquecidos. Não importa a diferença pela qual a primeira classificada venceu a segunda, só obtém uma vantagem de meio ponto. Depois, para a execução do programa longo, é usado o mesmo sistema de pontuação, com a única diferença de que os pontos são a dobrar, pelo que a primeira classificada recebe agora 1 ponto, a segunda 2 pontos, a terceira 3 pontos, etc. Em seguida, somam-se os pontos dos dois programas, produzindo a pontuação global de cada patinadora. O total mais baixo ganha a medalha de ouro. Depois de Hughes, Kwan e Cohen terem apresentado os seus programas longos, Hughes estava à frente e tinha uma

pontuação de 1 no programa longo. Kwan era a segunda classificada, com 2 pontos, e Cohen estava em terceiro lugar, com 3 pontos. Somando todas as pontuações, vemos que antes de Slutskaya executar o seu programa, as pontuações globais eram:

1.º Kwan (2.5), 2.º Hughes (3.0), 3.º Cohen (4.5)

Finalmente, chega a vez de Slutskaya apresentar o seu programa e fica em segundo lugar no programa longo, pelo que agora as pontuações globais do programa longo são:

Hughes (1.0), Slutskaya (2.0), Kwan (3.0), Cohen (4.0)

O resultado é extraordinário. A vencedora da competição é Hughes porque as pontuações finais são:

1.º Hughes (3.0), 2.º Slutskaya (3.0), 3.º Kwan (3.5), 4.º Cohen (5.5)

Hughes ficou à frente de Slutskaya porque, quando há um empate nas pontuações finais, o desempenho superior no programa longo é usado para desempatar. Mas o efeito das regras mal concebidas é claro. O desempenho de Slutskaya resulta na troca de posições entre Kwan e Hughes. Kwan está à frente de Hughes depois de ambas terem apresentado o seu programa, mas depois da apresentação de Slutskaya, Kwan acaba a ficar atrás de Hughes! Como é que os méritos relativos de Kwan e Hughes podem depender do desempenho de Slutskaya? É o paradoxo das alternativas irrelevantes.

A história é apenas uma coisa a seguir à outra.

Henry Ford

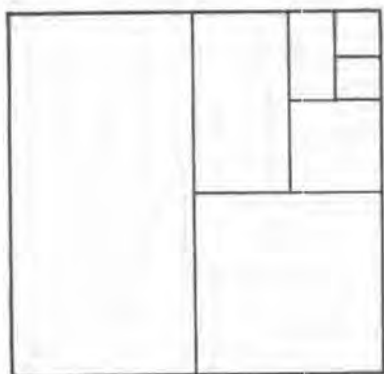
Os infinitos são uma coisa complicada e há milhares de anos que espantam os matemáticos e os filósofos. Às vezes, a soma de uma lista interminável de números torna-se infinitamente grande; outras vezes aproxima-se cada vez mais de um número determinado; às vezes resiste a qualquer tipo de soma definitiva. Há algum tempo, eu estava a dar uma palestra acerca do infinito que incluía a análise de uma simples série geométrica:

$$S = \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \frac{1}{16} + \frac{1}{32} + \frac{1}{64} \dots$$

E continuava, até ao infinito: cada termo da soma sendo exatamente metade do anterior. Na realidade, a soma desta série é igual a 1, mas um membro do público que não era um matemático quis saber se havia uma forma de lhe mostrar que isto era verdade.

Felizmente, há uma demonstração simples que usa apenas um desenho. Desenhe um quadrado com o tamanho 1 x 1, de modo a que a sua área seja 1. Agora divida o quadrado ao meio, dividindo-o na vertical em dois retângulos. Cada um destes retângulos deverá ter uma área igual a $\frac{1}{2}$. Em seguida, divida esses dois retângulos, de modo a formar dois retângulos mais pequenos, cada um com uma área de $\frac{1}{4}$. Em seguida divida esses dois retângulos ao meio, de modo a formar mais dois retângulos, cada um deles com uma área de um $\frac{1}{8}$. Continue a seguir este procedimento, sempre criando um retângulo com uma área igual a metade da área do retângulo anterior e analise a imagem que obtém. A área do quadrado original

foi simplesmente subdividida numa sequência interminável de regiões que o enchem completamente. A área total do quadrado é igual à soma das áreas das peças que deixámos intactas em cada uma das fases do processo de divisão e a soma das áreas dessas peças é simplesmente igual à nossa série S . Desta forma, a soma da série S tem de ser igual a 1, que é a área total do quadrado.



Normalmente, quando encontramos uma série como a S pela primeira vez, calculamos a sua soma de outra forma. Observamos que cada termo é igual a metade do termo anterior e multiplicamos a série toda por $1/2$ de modo a termos:

$$\frac{1}{2} \times S = \frac{1}{4} + \frac{1}{8} + \frac{1}{16} + \frac{1}{32} + \frac{1}{64} \dots$$

Mas observamos que a série à direita é apenas a série original, S , menos o primeiro termo, que é $\frac{1}{2}$. Assim, concluímos que $\frac{1}{2} \times S = S - \frac{1}{2}$ e novamente $S = 1$.

93 | SEGREGAÇÃO E “MICROMOTIVOS”

O mundo está cheio de coisas óbvias que nunca ninguém observa por acaso.

Sherlock Holmes, in *O Cão dos Baskerville*

Em muitas sociedades, dá-se uma segregação significativa entre comunidades de diferentes tipos – raciais, étnicas, religiosas, culturais e económicas. Em alguns casos, existe um desprezo manifesto de uma comunidade por parte de outra, mas noutros não parece haver alguma tentativa clara de separação e as diferentes comunidades convivem muito bem individualmente nas suas esferas de atividade. No entanto, as tendências dos indivíduos podem não ser um bom guia para o comportamento de um grupo por causa da interação entre muitas escolhas individuais. Quando alguns dos métodos estatísticos usados pelos cientistas para estudar o comportamento coletivo de grandes números de coisas são aplicados a populações de pessoas, emergem algumas verdades muito simples mas inesperadas.

Em 1978, Thomas Schelling, um cientista político americano, decidiu investigar como surgia a segregação racial nas cidades. Muitas pessoas supunham tratar-se apenas de intolerância racial. Outras pensavam que o problema podia ser resolvido se se criassem misturas aleatórias de diferentes comunidades, mas ficaram surpreendidas ao constatar que o resultado era sempre uma nova segregação em diferentes subcomunidades raciais, apesar de os residentes parecerem bastante tolerantes em relação aos outros grupos étnicos quando eram questionados. A conclusão que emergiu de um estudo matemático de sociedades virtuais que usava simulações informáticas foi que desequilíbrios muito ligeiros têm como consequência a segregação total, apesar da aparente tolerância. Suponhamos

que uma família muda de casa, por causa da intolerância ou para a evitar, se mais de 1 em cada 3 dos seus vizinhos for diferente deles, mas permanece na sua casa se menos de 1 em cada 5 forem diferentes. Nesta situação, uma combinação aleatória de dois tipos de famílias ("azuis" e "vermelhas") que sejam diferentes em algum aspeto (por exemplo, raça, religião ou classe) tornar-se-á cada vez mais polarizada até acabar por ser completamente segregada numa comunidade completamente azul ou completamente vermelha, como uma mistura de azeite e água*, separada por zonas neutras. Numa região com um número de vermelhos superior à média, os azuis optarão por ir embora, conduzindo à concentração de um número superior à média de azuis nos outros lugares, o que leva os vermelhos desses outros lugares a mudarem-se, e por aí em diante. As mudanças tendem a ser sempre para regiões onde há uma concentração acima da média de pessoas do tipo da que se mudou. As regiões fronteiriças que separam estas zonas são sempre zonas muito delicadas, uma vez que a mudança de uma única pessoa pode fazer a balança pender para um lado ou para o outro. Uma evolução mais estável seria no sentido de deixar estas regiões vazias de modo a criar uma zona neutra entre as comunidades segregadas.

Estas noções simples são muito importantes. Mostram que uma segregação muito forte é praticamente inevitável em comunidades mistas e não implica que exista um problema grave de intolerância. A segregação não implica necessariamente preconceito – embora possa, certamente, implicá-lo, como se pode constatar pelos exemplos dos Estados Unidos, do Zimbabué, da África do Sul e da Jugoslávia. É melhor promover a ligação das comunidades separadas do que tentar impedir a sua formação. O "macrocomportamento" é produzido por "micromotivos" que não têm de fazer parte de uma política organizada.

* Na realidade, este exemplo tão familiar não deve ser entendido à letra. Se o ar dissolvido for removido da água, por exemplo por uma sucessão de processos de congelação e descongelação, misturar-se-á facilmente com o azeite.

O e-mail é uma coisa maravilhosa para aquelas pessoas cuja função na vida é estar no topo, a controlar tudo. Não é para mim: a minha função é estar na base, na origem das coisas.

Don Knuth

No capítulo anterior vimos um exemplo de comportamento coletivo em que nenhum indivíduo se quer encontrar numa minoria significativa. Nem todas as situações são assim. Se está a pensar ir de férias para uma ilha paradisíaca, quer pertencer à minoria e não à maioria, no que diz respeito ao destino de férias escolhido pela maioria das pessoas. O bar que “toda a gente” escolhe frequentar por causa da música ou da comida vai acabar por se tornar uma experiência muito pouco ideal se tiver de ficar na fila de espera para entrar, não encontrar uma mesa vaga ou tiver de esperar uma hora para ser atendido. Será mais feliz num bar menos popular.

É como um jogo em que “ganha” por pertencer à minoria. Habitualmente há um número médio de pessoas que escolhem ir a todos os locais públicos disponíveis, mas as flutuações em volta da média são muito significativas. Para as reduzir e convergir numa estratégia mais útil, é necessário usar as informações anteriores acerca das pessoas que frequentam esse local. Se simplesmente tentar adivinhar a psicologia dos outros clientes, acabará a cometer o erro habitual de pensar que não pertence à média. Pensará que a sua escolha não será também feita por muitas outras pessoas que estão a agir com base nas mesmas evidências – e é assim que descobre que todas as outras pessoas também decidiram ir passear à beira do rio numa tarde de domingo soalheira.

Se houver dois locais à escolha, como resultado da experiência acumulada de todas as pessoas, a estratégia ideal aproxima-se cada vez mais de metade das pessoas escolher cada um dos locais – pelo que nenhum

deles é especialmente popular ou impopular – em média. A princípio, as flutuações em torno da média são bastante grandes e poderá chegar a um local para descobrir que tem menos clientes do que é habitual. À medida que o tempo passa, usará cada vez mais as suas experiências passadas para determinar quando e se estas flutuações irão ocorrer e agirá em conformidade, tentando ir ao local com menos pessoas. Se todas as pessoas se comportarem da mesma forma, o local manterá o mesmo número de clientes ao longo do tempo, mas as flutuações diminuirão a um ritmo constante. O último ingrediente desta situação é que haverá pessoas que confiam na sua memória e nas suas análises das experiências passadas, e outras que não o fazem ou que só recorrem à experiência algumas vezes e que têm de fazer uma escolha. Isto tende a dividir a população em dois grupos – os que seguem completamente as experiências passadas e os que as ignoram. Uma vez que as consequências de fazer a escolha errada são muito mais negativas (não jantar ou ter a noite arruinada) do que as consequências de fazer a escolha certa (jantar mais depressa e ter uma noite mais agradável), as pessoas têm mais cuidado para evitarem fazer a escolha errada e tendem a distribuir igualmente as suas apostas pelas duas escolhas disponíveis, tendo uma probabilidade igual de acertar a longo prazo. Adotar uma estratégia mais arriscada resulta numa quantidade mais extrema de perdas do que de ganhos. O resultado é um padrão muito cauteloso de decisão de grupo que está longe de ser ideal e todos os restaurantes ficam vazios.

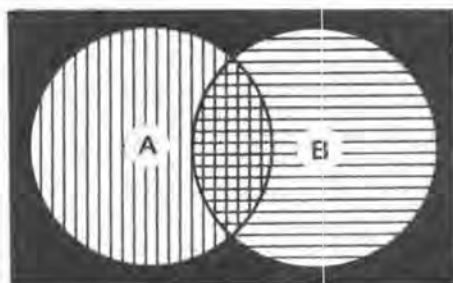
Há dois grupos de pessoas no mundo; as que acreditam que o mundo pode ser dividido em dois grupos de pessoas e as que não acreditam.

Anônimo

Jon Venn era oriundo da zona este de Inglaterra, de perto da vila piscatória de Hull e foi – como acontecia com os matemáticos promissores – para Cambridge, onde entrou na Gonville and Caius College em 1853. Tendo acabado o curso de matemática entre os seis melhores, foi selecionado para um estágio como professor universitário. Depois deixou a faculdade durante quatro anos e foi ordenado padre em 1859, seguindo os passos do pai e do avô, que eram figuras eminentes na ala evangélica da Igreja Anglicana. No entanto, em vez de seguir o caminho eclesiástico que tinha sido aberto para si, regressou a Caius em 1862 para dar aulas de lógica e probabilidades. Apesar desta estranha combinação de acaso, lógica e teologia, Venn também era um homem prático e muito bom a construir máquinas. Construiu uma máquina de lançamento de bolas de críquete suficientemente boa para derrotar um dos membros da equipa australiana em quatro ocasiões, quando a equipa visitou Cambridge em 1909.

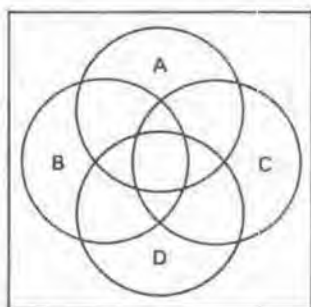
Foram as suas palestras sobre lógica e probabilidades na universidade que o tornaram famoso. Em 1880, apresentou um diagrama para representar possibilidades lógicas. Este rapidamente substituiu as alternativas que tinham sido experimentadas pelo grande matemático suíço Leonard Euler e pelo lógico de Oxford e escritor surrealista vitoriano Lewis Carroll. Acabou por ser designado “Diagrama de Venn” em 1918.

Os diagramas de Venn representavam possibilidades usando regiões de espaço. Aqui está um exemplo simples destes diagramas que representa todas as possibilidades quando há dois atributos:



Suponhamos que A é um conjunto de todos os animais castanhos e que B é um conjunto de todos os gatos. A região quadriculada onde há sobreposição contém todos os gatos castanhos; a parte da região A que não se cruza com a região B contém todos os animais castanhos que não são gatos; a parte da região B que não se cruza com a região A representa todos os gatos que não são castanhos; e, finalmente, a região a negro, fora das esferas A e B, representa tudo o que não é nem um animal castanho nem um gato.

Estes diagramas são muito usados para apresentar todos os diferentes conjuntos de possibilidades que podem existir. Estão limitados pela "lógica" da página bidimensional em que são desenhados. Suponhamos que representamos quatro conjuntos, usando os círculos A, B, C e D. Estes vão representar as ligações de amizade entre três pessoas que podem existir em grupos de quatro pessoas: Alex, Bob, Chris e Dave. A região A representa as amizades mútuas entre Alex, Bob e Chris; a região B as amizade entre Alex, Bob e Dave; a região C representa as amizades entre Bob, Chris e Dave; e a região D representa as amizades entre Chris, Dave e Alex. A forma como o diagrama de Venn foi desenhado apresenta uma sub-região em que A, B, C e D se cruzam. Isto significaria que a região onde os quatro círculos se sobrepõem contém alguém que é membro de A, B, C e D. No entanto, essa pessoa que é comum aos quatro conjuntos não existe.



96 | ALGUNS BENEFÍCIOS DA IRRACIONALIDADE

O misticismo pode ser caracterizado como um estudo das proposições que são equivalentes às suas próprias refutações. Segundo o ponto de vista ocidental, a classe de todas essas proposições está vazia. Segundo o ponto de vista oriental, tal classe só pode estar vazia se não existir.

Raymond Smullyan

As fotocópias não são tão simples como parecem. Se tirar uma aqui na Europa, rapidamente apreciará uma característica a que está tão habituado que já a toma por certa. Se colocar duas folhas A4 na fotocopadora, com as faces voltadas para baixo, poderá reduzi-las de modo a serem impressas lado a lado numa folha de papel A4. A redução fica exatamente à medida da folha de papel e não há margens irregulares na página da cópia final. No entanto, se experimentar fazê-lo com uma folha de papel de carta com o padrão americano, obterá um resultado muito diferente. O que é que acontece neste caso e o que é que isso tem que ver com a matemática e com a irracionalidade?

A normalização internacional dos formatos de papel (ISO), da qual o A4 é um formato possível, deriva de uma observação simples inicialmente feita pelo físico alemão Georg Lichtenberg, em 1786. Cada um dos formatos de papel da série A tinha metade da área da folha do tamanho acima por ter metade da largura, mas o mesmo comprimento. Assim, se colocarmos duas folhas lado a lado, criamos uma folha do tamanho acima: por exemplo, duas folhas A4 fazem uma folha A3. Se o comprimento for C e a largura for L , isto significa que têm de ser escolhidas de modo a que $C/L = 2L/C$. Isto exige que $C^2 = 2L^2$, pelo que os comprimentos dos lados estejam em proporção com a raiz quadrada de 2, um número irracional que é aproximadamente igual a 1.41: $C/L = \sqrt{2}$.

Esta proporção irracional entre o comprimento e a largura de cada formato de papel, chamada “relação de aspeto” do papel, é a característica que define a série de formatos A. A maior folha, chamada A0, é definida como tendo uma área de 1 metro quadrado, pelo que as suas dimensões são $C(A0) = 2^{1/4}$ m e $L(A0) = 2^{-1/4}$ m, respetivamente. A relação de aspeto significa que uma folha de tamanho A1 tem um comprimento de $2^{-1/4}$ e uma largura de $2^{-3/4}$, pelo que a sua área é de apenas $\frac{1}{2}$ metro quadrado. Continuando com este padrão, poderá querer considerar as dimensões de uma folha de papel AN, em que $N = 0, 1, 2, 3, 4, 5, \dots$, etc., será:

$$C(AN) = 2^{1/4+N/2} \text{ e } L(AN) = 2^{-1/4+N/2}$$

A área de uma folha de papel com o formato AN será, desta forma, igual à largura vezes o comprimento, que é igual a 2^{-N} metros quadrados.

Podiam ter sido escolhidos vários tipos de relações de aspeto em vez da $\sqrt{2}$. Se preferisse, poderia ter optado pela Proporção Dourada, tão apreciada pelos artistas e arquitetos dos tempos antigos. Esta escolha corresponderia a uma escolha de formatos de papel com $C/L = (C+L)/C$, em que $C/L = (1+\sqrt{5})/2$, mas na prática esta não seria uma escolha sensata.

A beleza da relação de aspeto $\sqrt{2}$ torna-se mais óbvia se regressarmos ao exemplo da fotocopiadora. Significa que se pode reduzir o tamanho de um A3, ou de duas folhas A4 postas lado a lado, para caber numa única folha A4 sem deixar um espaço em branco na cópia final. Poderá observar que a maioria das fotocopiadoras oferece a possibilidade de redução de 70% (ou 71% se for uma marca mais preciosa) do A3 para o A4. A razão de ser disto é que 0,71 é aproximadamente igual a $1/\sqrt{2}$ e é a medida certa para reduzir uma folha A3 ou duas folhas A4 para caberem numa folha A4. Duas dimensões, C e L, são reduzidas para $C\sqrt{2}$ e $L\sqrt{2}$, de modo a que a área de CL seja reduzida para $CL\sqrt{2}$, conforme exigido para reduzir uma folha de qualquer tamanho AN para o tamanho abaixo. Da mesma forma, para ampliações, o número apresentado no painel de controlo da fotocopiadora é 140% (ou 141% em algumas fotocopiadoras) porque $\sqrt{2} = 1,41$, aproximadamente. Outra consequência desta relação de aspeto constante para todas as reduções e ampliações é que os diagramas mantêm as mesmas formas relativas: os quadrados não se transformam em retân-

gulos e os círculos não se transformam em elipses quando o seu tamanho é mudado de um formato para outro da série A.

De um modo geral, as coisas são diferentes na América e no Canadá. Os formatos de papel determinados pelo American Standards Institute (ANSI) que são usados nestes países, em polegadas pois foi assim que foram definidos, são A ou Letter (8,5 pol. x 11,0 pol.), B ou Legal (22 pol. x 17 pol.), C ou Executive (17 pol. x 22 pol.), D Ledger (22 pol. x 34 pol.) e E Ledger (34 pol. x 44 pol.). Têm duas relações de aspeto distintas: 17/11 e 22/17. Assim, se quiser manter a mesma relação de aspeto quando une dois formatos de papel, precisa de saltar dois tamanhos em vez de apenas um. Consequentemente, não é possível ampliar ou reduzir duas folhas de papel para uma folha do tamanho abaixo sem deixar uma margem em branco na folha. Para fazer cópias ampliadas ou reduzidas numa fotocopiadora americana é preciso mudar o tabuleiro do papel de modo a aceitar papéis com duas relações de aspeto em vez de usar o fator $\sqrt{2}$ único que usamos no resto do mundo. Às vezes, um pouco de irracionalidade é útil.

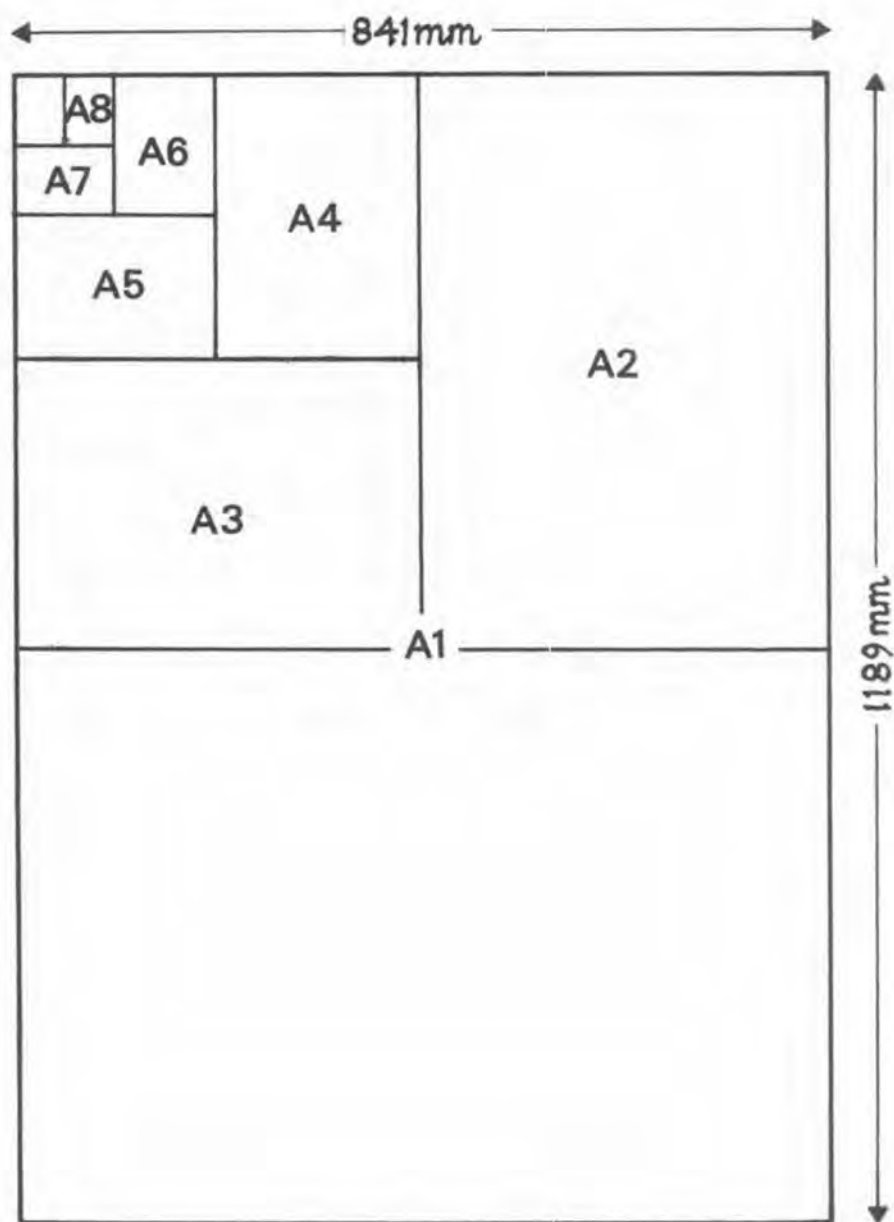


Diagrama dos formatos de papel da série A segundo a norma ISO

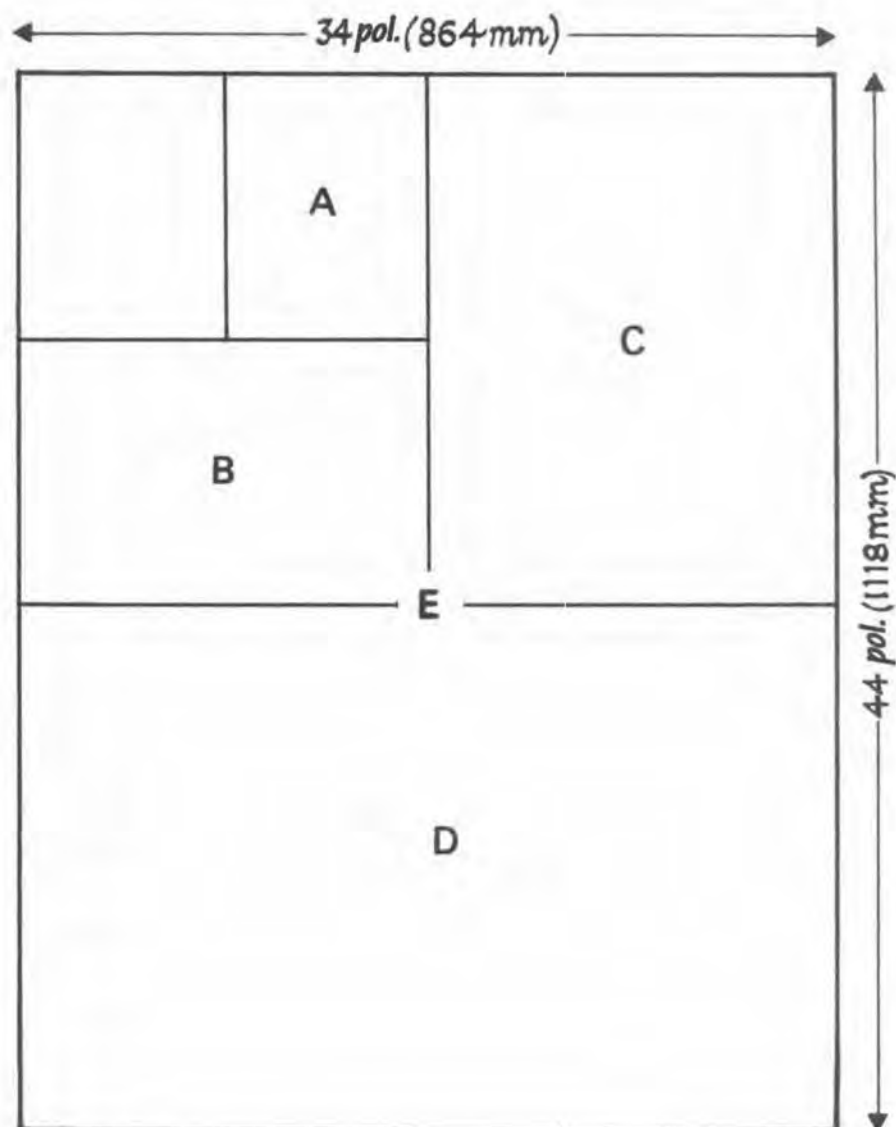


Diagrama dos formatos de papel segundo a norma ANSI

Decisão mais ação vezes planeamento é igual a produtividade menos atraso ao quadrado.

Armando Iannucci

A matemática tornou-se de tal forma um símbolo de *status* em alguns meios que há uma necessidade de a usar sem pensar se é adequado. Só porque pode usar símbolos para expressar uma palavra, não significa necessariamente que isso aumente o seu conhecimento. Dizer “Três Porquinhos” é mais útil do que definir o conjunto de todos os porcos, o conjunto de todos os grupos de três e o conjunto de todos os animais pequenos e isolar a parte em que os três conjuntos se sobrepõem. Um interessante passo nesta direção foi dado pela primeira vez em 1725 pelo filósofo escocês Francis Hutcheson, que veio a tornar-se um professor de filosofia muito bem sucedido na Universidade de Glasgow graças a esse seu trabalho. Ele queria calcular a bondade moral de ações individuais. Vemos aqui parte do impacto do sucesso de Newton ao descrever o mundo físico através da matemática: a sua metodologia foi copiada e admirada em todos os domínios. Hutcheson propôs uma fórmula universal para avaliar a virtude, ou grau de benevolência, das nossas ações:

$$\text{Virtude} = \frac{\text{Bem Público} \pm \text{Interesse Privado}}{\text{Capacidade Natural para Fazer o Bem}}$$

A fórmula de Hutcheson para a aritmética moral tem várias características agradáveis. Se duas pessoas tiverem a mesma capacidade natural para fazer o Bem, a que produzir a maior quantidade de bem público é mais virtuosa. Da mesma forma, se duas pessoas produzirem o mesmo

nível de bem público, a que tiver menor capacidade natural para o fazer é mais virtuosa.

O outro ingrediente da fórmula de Hutcheson, o Interesse Privado, pode contribuir positiva ou negativamente ($\pm$). Se a ação de uma pessoa beneficiar o público, mas a prejudicar a ela (por exemplo, se fizer trabalho de caridade sem ser paga em vez de aceitar um emprego remunerado), a Virtude é aumentada pelo Bem Público + Interesse Privado. Mas se as ações ajudarem o público e a pessoa que as pratica (por exemplo, fazer campanha para pôr fim a um projeto de construção que prejudica tanto a sua casa como as dos seus vizinhos), a Virtude da ação é diminuída pelo fator Bem Público – Interesse Privado.

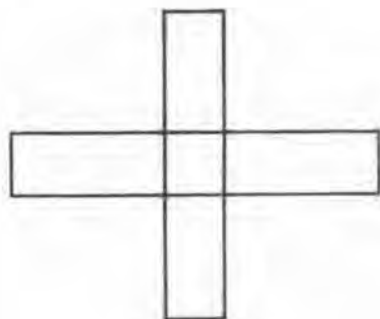
Hutcheson não atribuiu valores numéricos às quantidades na sua fórmula, mas estava preparado para os adotar se fosse necessário. A fórmula moral não é realmente útil porque não revela nada de novo. Todas as informações fazem parte da sua própria estrutura. Qualquer tentativa de calibrar as unidades da Virtude, do Interesse Próprio e da Capacidade Natural seria inteiramente subjetiva e nunca poderia ser feita uma previsão mensurável. Ainda assim, a fórmula é uma abreviatura muito prática para muitas palavras.

Algo estranhamente remanescente da fantasia racionalista de Hutcheson apareceu 200 anos mais tarde, num projeto do famoso matemático americano George Birkhoff, que ficou intrigado com o problema de quantificar a apreciação estética. Dedicou um longo período da sua carreira à busca de uma forma de quantificar aquilo que nos agrada na música, na arte e no *design*. Os seus estudos recolheram exemplos de muitas culturas e o seu livro *Aesthetic Measure* ainda hoje é fascinante. Notavelmente, reduz tudo a uma única fórmula que me faz lembrar Hutcheson. Ele acreditava que as qualidades estéticas são determinadas por uma medida que é determinada pela proporção de ordem e complexidade:

$$\text{Medida Estética} = \text{Ordem/Complexidade}$$

Procurou criar formas de calcular a Ordem e Complexidade de padrões e formas específicos de um modo objetivo e aplica esses critérios a todo o tipo de formas de vasos, padrões de azulejos, frisos e desenhos. Claro que, como em qualquer avaliação estética, não faz sentido comparar vasos

de cerâmica com quadros: temos de nos cingir a um meio e forma específico, para que este raciocínio possa fazer sentido. No caso das formas poligonais, a medição da Ordem por Birkhoff soma pontos pela presença ou ausência de quatro possíveis simetrias diferentes e subtrai uma penalização (de 1 ou 2) para certos ingredientes insatisfatórios (por exemplo, se as distâncias entre os vértices forem demasiado pequenas ou se os ângulos interiores estiverem demasiado próximos de 0 ou de 180 graus, ou se houver uma falta de simetria). O resultado é um número que nunca pode ser superior a 7. A Complexidade é definida como o número de linhas retas que contém um polígono. Assim, para um quadrado a classificação é 4, mas para uma cruz romana (como aquela aqui apresentada) é 8 (4 horizontais mais 4 verticais):



A fórmula de Birkhoff tem o mérito de usar números para classificar os elementos estéticos, mas infelizmente a complexidade estética é demasiado vasta para poder ser englobada por uma forma tão simples e, ao contrário da tentativa mais rudimentar de Hutcheson, não cria uma medida com a qual muitas pessoas concordem. Se aplicarmos a fórmula dele aos padrões fractais modernos que são apreciados por tantas pessoas (e não apenas matemáticos), com os seus padrões repetitivos em escalas cada vez menores, a Ordem destes pode obter uma pontuação superior a 7, mas a sua Complexidade torna-se progressivamente maior à medida que o padrão explora escalas cada vez menores e a Medida Estética tende rapidamente para o zero.

Os outros países têm futuros imprevisíveis, mas a Rússia é um país com um passado imprevisível.

Yuri Afanasiev

O caos é a sensibilidade extrema à ignorância. Surge em situações em que um pouco de ignorância acerca da situação atual cresce gradualmente com o passar do tempo, não aumentando apenas em proporção com o tempo decorrido, mas duplicando a cada passo. O exemplo mais famoso deste tipo de situação é o tempo. Não são raras as vezes em que as previsões do tempo são menos do que acertadas, mas isso não significa que não somos bons no que fazemos ou que há algum segredo oculto da meteorologia que os físicos ainda não descobriram – deve-se ao nosso conhecimento imperfeito do estado atual do tempo. Temos estações meteorológicas a cada 100 km espalhadas por todo o país, e menos estações destas no mar, que efetuam medições periódicas. No entanto, isto continua a deixar margem para variações consideráveis entre estações. Infelizmente, pequenas diferenças nessa extrapolação conduzem frequentemente a condições meteorológicas futuras muito diferentes.

Este tipo de sensibilidade do futuro ao presente começou a ser estudado exaustivamente na década de 1970, quando os computadores domésticos se tornaram mais acessíveis e os cientistas passaram a ter um acesso mais facilitado a eles. Foi designado por caos para refletir os resultados inesperados que podiam ser obtidos em pouco tempo a partir de condições iniciais aparentemente inócuas, devido ao rápido crescimento dos efeitos de qualquer nível de incerteza, por muito pequeno que fosse. A indústria cinematográfica aproveitou a ideia para o filme *Parque Jurássico*, em que um pequeno erro, que levou à reprodução cruzada de dinossauros, e um tubo de ensaio partido produziram um grande desastre; as coisas foram

de mal a pior e um pequeno nível de incerteza tornou-se rapidamente em ignorância quase total. Até havia um especialista do caos disponível para nos explicar tudo enquanto víamos os problemas agravarem-se ao ponto de fugirem completamente de qualquer possibilidade de controlo.

Há um aspeto muito interessante do caos que está associado às experiências não matemáticas que obtemos a partir dos livros, da música e do teatro. Como é que determinamos se um livro, uma peça de teatro ou uma música é “bom” ou melhor do que qualquer outro? Por que é que consideramos que *A Tempestade* é melhor do que *A Espera de Godot*, ou a *Quinta Sinfonia* de Beethoven superior ao 4’ 33” de John Cage, que é composto por 4 minutos e 33 segundos de silêncio?

Um argumento popular diz que os bons livros são aqueles que queremos voltar a ler, as boas peças deixam-nos com vontade de voltar a vê-las e a boa música produz um desejo de ver outro concerto. Procuramos fazer estas coisas porque estas obras têm uma certa quantidade de imprevisibilidade caótica. Uma pequena alteração na encenação ou no elenco da peça *A Tempestade*, uma orquestra e um maestro diferentes num concerto de música clássica, por exemplo, resultariam numa experiência muito diferente da peça ou da música.

A arte banal não tem essa qualidade. Uma mudança das circunstâncias produziria basicamente a mesma experiência. Não há necessidade de repetir a experiência. O caos não é apenas algo que devemos evitar ou controlar.

Há quem pense que a possibilidade do caos é o fim da ciência. Tem de haver sempre um determinado nível de ignorância em relação a tudo no mundo – não temos instrumentos perfeitos. Como podemos esperar prever ou compreender algo se houver um crescimento rápido destas incertezas? Felizmente, embora os átomos e moléculas individuais se movam de forma caótica na sala onde estou, o seu movimento, em média, é completamente previsível. Muitos sistemas caóticos são assim e usamos realmente algumas dessas quantidades médias para medir o que está a acontecer. A temperatura, por exemplo, é uma medida da velocidade média das moléculas na sala. Apesar de as moléculas individuais terem uma história impossível de prever ao fim de algumas colisões com as moléculas vizinhas e com outros objetos mais densos, estas colisões mantêm a média regular e previsível. O caos não é o fim do mundo.

* Surpreendo-me sempre quando descubro que a minha “revelação” do motivo por que são quatro minutos e trinta e três segundos (e não outro número qualquer) é novidade para todos os músicos com quem falo. Na realidade, o intervalo de 273 segundos foi escolhido por Cage para o seu Zero Absoluto de som em analogia com os 273 graus Celsius negativos que são o Zero Absoluto ou a temperatura a que todos os movimentos moleculares clássicos param. Incrivelmente, Cage afirmou que esta foi a sua obra mais importante.

Não sou um apaixonado das alturas, não sou um arossexual. Não gosto de aviões. Nunca quis ser piloto como todos aqueles bandos de rufias que povoam a indústria aeronáutica.

Michael O'Leary, dono da Ryanair

Se já passou tanto tempo como eu em filas para embarcar em aviões, conhece todos os erros que podem acontecer. As companhias aéreas *low-cost*, como a Ryanair, simplesmente não querem saber: é um vale-tudo sem lugares reservados. Perceberam que havia um incentivo para tornar o serviço o pior possível durante algum tempo, para depois poderem vender-nos o direito de “embarque prioritário”, à frente de todos os outros passageiros. Não há prioridade para os acompanhantes de crianças de colo nem para as pessoas com problemas de mobilidade, pelo que estes passageiros tornam o embarque ainda mais lento. E o que acontece quando toda a gente escolhe o embarque prioritário? Não sei, mas desconfio que esse é o objetivo final.

As companhias aéreas tradicionais têm uma variedade de métodos para reduzir o stresse e os atrasos para os passageiros da classe turística. Todos os passageiros têm um lugar atribuído e as crianças e pessoas com problemas de mobilidade são as primeiras a embarcar. Em algumas companhias, o embarque é feito por ordem de lugares, pelo que quando só há uma entrada na parte da frente, as pessoas cujos lugares ficam na parte de trás do avião embarcam primeiro, para não obstruírem a passagem aos outros passageiros. No papel isto parece tudo muito bonito, mas na prática há sempre alguém que bloqueia a passagem ao tentar enfiar uma mala demasiado grande nos compartimentos de bagagem de mão, as pessoas nos lugares de coxía têm de se levantar constantemente para deixar passar as que têm lugares à janela, e toda a gente atrapalha toda a gente. Tem de haver um sistema melhor.

Um jovem astrofísico alemão, Jason Streffen, que trabalha no Fermi Lab, perto de Chicago, teve a mesma ideia e começou a explorar a eficiência de diferentes estratégias de embarque de passageiros usando uma simulação informática simples que permitia fazer modificações à estratégia de embarque e adicionar todas as variações aleatórias que se quisesse para perturbar mesmo os planos mais bem concebidos. O avião virtual que criou tinha 120 lugares, levava 6 passageiros por fila, era dividido ao centro por uma coxia e não tinha classe executiva nem primeira classe. Todos os passageiros virtuais transportavam bagagem de mão.

Foi fácil encontrar a pior política de embarque a adotar para um avião com uma única entrada na parte da frente: fazer embarcar os passageiros de acordo com os números dos lugares, a começar pela parte da frente. Todos os passageiros teriam de lutar para passar pelos que já estavam dentro do avião e que estavam ocupados a colocar a bagagem de mão nos compartimentos, para poderem chegar ao seu lugar. Isto é bastante óbvio, mas levou as companhias aéreas a concluir que a melhor estratégia era simplesmente o contrário da pior: fazer embarcar os passageiros por ordem de lugares a começar pelo fundo do avião. Curiosamente, a investigação de Streffen mostrou que este era apenas o segundo método de embarque mais rápido! O único método pior era começar por embarcar os passageiros cujos lugares ficavam na parte da frente. Até uma ordem completamente aleatória de embarque, independentemente do número do lugar, obtinha resultados muito melhores. Mas o melhor método de todos era estruturado. Os passageiros deviam embarcar de modo a que os lugares junto à janela fossem ocupados antes dos lugares do meio ou da coxia, e deviam embarcar por uma ordem que distribuisse as pessoas que estavam a arrumar a bagagem de mão por todo o comprimento do avião em vez de se concentrarem todas numa única zona.

Se os passageiros em todas as filas com números pares e com lugares à janela embarcassem primeiro, teriam um espaço desimpedido na coxia à sua frente e atrás de si e não impediriam a passagem às outras pessoas enquanto estavam a arrumar a bagagem. Todas as pessoas poderiam arrumar a bagagem de mão ao mesmo tempo. Se alguém precisasse de passar, haveria espaços adicionais na coxia para se desviarem. Começando pelo fundo do avião, a necessidade de passar pelos outros passageiros é mais uma vez evitada. A seguir, embarcariam os passageiros com os lugares

do meio e da coxia. Depois embarcariam os passageiros das filas com números ímpares.

Nem todos terão facilidade em seguir esta estratégia à letra. As crianças pequenas terão de permanecer junto dos pais, mas podem embarcar primeiro ainda assim. No entanto, o tempo ganho com a implementação desta estratégia pode ser considerável. A simulação informática mostrou que ao longo de centenas de experiências com pequenas variações ("passageiros problemáticos"), este método revelou-se, em média, sete vezes mais rápido do que o método de fazer embarcar os passageiros segundo o sistema tradicional, a começar pelo fundo do avião. Streffen registou a patente deste sistema!

Imagem todas as pessoas
A partilhar o mundo inteiro.

John Lennon, "Imagine"

As vezes não conseguimos ver o bosque no meio das árvores. Os números muito grandes toldam-nos a compreensão. É difícil imaginar um milhão, quanto mais mil milhões. Reduzir o tamanho das coisas ajuda-nos a torná-las mais concretas e imediatas. Em 1900, foi sugerida a famosa ilustração do estado do mundo*, que nos pedia que imaginássemos que a população do mundo era reduzida em escala de modo a criar uma única aldeia de 100 habitantes, com todos os atributos reduzidos à mesma escala. Que aspeto teria essa aldeia?

Conteria 57 Asiáticos e 21 Europeus, 14 membros do hemisfério ocidental, tanto norte como sul, e apenas 8 Africanos. Dos 100 aldeões, 70 seriam de uma raça que não a branca e 30 seriam de raça branca. Apenas 6 pessoas teriam 59% da riqueza da aldeia e essas 6 pessoas seriam dos Estados Unidos. Por fim, de entre os 100 habitantes da aldeia, 80 viveriam em casas de má qualidade, 70 seriam analfabetos, 50 sofreriam de subnutrição, apenas 1 teria um computador e 1 teria formação superior. Que aldeia tão estranha, não acha?

* Esta imagem da aldeia global foi proposta inicialmente por Donatella Meadows e baseava-se numa população de 1000 pessoas. Depois de ouvir Meadows a falar na rádio, o ativista pelo ambiente David Copeland localizou Meadows e reviu a estatística dela para refletir uma população de 100 pessoas. Acrescentou a imagem a um cartaz que ia ser distribuído a 50 000 pessoas na Cimeira da Terra no Rio de Janeiro. O exemplo de Meadows do relatório do Estado da Aldeia foi publicado em 1990 com o título *Who Lives in the Global Village?* Desenvolveu-se a partir daqui o mito de que este exemplo foi originado por um professor da universidade de Stanford, Philip Harter, mas na realidade ele limitou-se a reencaminhar um e-mail com os factos compilados por Meadows e Copeland. Consulte o website http://www.odtmaps.com/behind_the_maps/population_map/state-of-villagestat.asp

Mergulhar num livro ou num artigo extenso costumava ser fácil para mim. A minha mente ficava presa à narrativa ou às voltas do enredo e passava horas a divagar por longas extensões de prosa. Agora isso raramente acontece. A minha concentração começa a divagar ao fim de apenas duas ou três páginas. Fico inquieto, perco o fio à meada e começo à procura de outra coisa para fazer. Sintomo como se estivesse constantemente a forçar a minha atenção a voltar para o texto. A leitura profunda que costumava ser tão natural tornou-se uma luta.

Nicholas Carr

¹ Presume-se que a corrente tem uma massa uniforme por unidade de comprimento, é perfeitamente flexível e tem uma espessura nula. A função matemática $\cosh(\dots)$ é definida em termos da função exponencial por $\cosh(x) = (e^x + e^{-x})/2$.

² Suponhamos que a aceleração máxima é $+A$, sendo a desaceleração máxima $-A$, e que tem de empurrar o carro da posição $x = 0$ para $x = 2D$. Começa a uma velocidade zero e acaba a uma velocidade zero. Se aplicar $+A$ à distância de $X = 0$ para $x = D$, será necessário um período de $\sqrt{(2D/A)}$ para chegar a $x = D$ e alcançará este ponto com uma velocidade $\sqrt{(2DA)}$. Se aplicar automaticamente a desaceleração $-A$ ao carro, este regressará à velocidade zero a $x = 2D$ depois de o período de desaceleração ter tido uma duração de $\sqrt{(2D/A)}$. Conforme é esperado devido à simetria da situação, este é exatamente o mesmo tempo que foi necessário para chegar a meio do caminho. Assim, o tempo total necessário para levar o carro para dentro da garagem é $2\sqrt{(2D/A)}$.

³ Os passos aleatórios são um processo de difusão governado pela equação da difusão. A difusão de y num tempo, t , e numa dimensão de distância, x , traduz-se em $\partial y / \partial t = K \partial^2 y / \partial x^2$ em que K é uma medida constante da facilidade de difusão através do meio. Portanto, se inspecionarmos as dimensões, concluímos que y/t deve ser proporcional a y/x^2 . Portanto, t é proporcional a x^2 e serão necessários P^2 passos para percorrer a distância em linha reta $x = P$.

⁴ Este comprimento é conhecido na física como comprimento Planck, em homenagem a Max Planck, um dos pioneiros da teoria quântica que foi o primeiro a determiná-lo. É a única quantidade com as dimensões de um comprimento que pode ser formada a partir das três grandes constantes da natureza, c , a velocidade da luz; h , a constante quântica de Planck; e G , a constante da gravitação. É igual a $(\hbar h/c^3)^{1/2}$ e é um reflexo singular do carácter relativista, quântico e gravitacional do universo. É uma unidade de comprimento que não é seleccionada pela conveniência humana e é por isso que parece tão pequena em comparação com as unidades que usamos na nossa vida quotidiana.

⁵ Um exemplo simples é o puzzle dos macacos em que temos, por exemplo, 25 peças quadradas, cada uma com 4 lados. Em cada um dos lados há metade de um macaco (a parte de cima ou a parte de baixo). Os macacos são apresentados em quatro cores e é preciso unir as 25 peças de modo a formar um quadrado maior, de 5×5 , de modo a que cada um dos lados complete a imagem do corpo de um macaco ao unir-se à peça adjacente com a mesma cor. Quantas formas alternativas de dispor as peças precisaria de analisar para encontrar a resposta "correta" que lhe permitiria resolver o puzzle, resultando numa imagem completa, em que todas as metades dos macacos foram unidas à outra metade com a mesma cor? Há 25 formas de colocar a primeira peça, seguida de 24 formas de colocar a segunda peça, 23 para a seguinte, etc. Portanto são $25 \times 24 \times 23 \times 22 \times \dots \times 3 \times 2 \times 1 = 25!$ formas de combinar as diferentes peças. Mas há 4 orientações possíveis para cada peça, o que cria mais 4^{25} possibilidades. O número total de configurações que será necessário experimentar para encontrar o padrão correto é $25! \times 4^{25}$. É um número incrivelmente grande. Se tentássemos escrevê-lo por extenso, não caberia na página. Se um computador tentasse examinar este enorme número de combinações a uma velocidade de um milhão por segundo, demoraria ainda assim mais de 5333

milhões de milhões de milhões de milhões de milhões de milhões de anos a examinar todos os números para encontrar a resposta certa. Para ter um termo de comparação, o universo está a expandir-se lá apenas 13,7 milhões de milhões de anos.

* É de notar que o melhor candidato é o que fica no $(r+1)$ º lugar e que estamos a deixar passar os primeiros r candidatos, pelo que iremos certamente escolher o melhor candidato, mas essa situação ocorrerá apenas com uma probabilidade de $1/N$. Se o melhor candidato estiver no $(r+2)$ º lugar, selecioná-lo-emos com uma probabilidade $1/N \times (r/(r+1))$. Se continuarmos até aos lugares mais afastados da nossa lista, veremos que a probabilidade de sucesso geral é apenas a soma dessas quantidades, que é $P(N,r) = 1/N \times [1 + r/(r+1) + r/(r+2) + r/(r+3) + r/(r+4) + r/(r+5) + \dots + r/(N-1)] \approx 1/N \times [1 + r \ln((N-1)/r)]$. Esta última quantidade na direção da qual a série converge quando N se torna muito grande, alcança o seu valor máximo quando o logaritmo $\ln[(N-1)/r] = 1$, pelo que $r = (N-1)/e \approx N/e$ em que N é um número grande. Desta forma, o valor máximo de $P(N,r)$ que ocorre é $P \approx r/N \times \ln(N/r) \approx 1/e \approx 0,37$.

7 Mais precisamente, se deixarmos passar os primeiros N/e candidatos, a nossa probabilidade de selecionar o melhor aproxima-se de $1/e$ quando N se torna um número muito grande, em que $e = 2,71828\dots = 1/0,37$ é a constante matemática que define a base dos logaritmos naturais.

* A probabilidade de uma pessoa não fazer anos no mesmo dia que o leitor é de $364/365$, portanto, se tiver C convidados e as datas de aniversário deles forem independentes, a probabilidade de nenhum deles partilhar a sua data de aniversário é $P = (364/365)^C$. Desta forma, a probabilidade de um dos convidados fazer anos no mesmo dia que o leitor é de $1 - P$. À medida que o número C se torna cada vez maior, o número P aproxima-se de zero e a probabilidade de ter a mesma data de aniversário que um dos seus convidados aproxima-se de 1, e verificamos que $1 - P$ ultrapassa os 0,5 quando C excede $\log(0,5)/\log(364/365)$ que é aproximadamente 253.

* Mais uma vez, é mais fácil começar por calcular a probabilidade de não terem a mesma data de aniversário. Se tiver N convidados, esta probabilidade é igual a $P = 365/365 \times 364/365 \times 363/365 \times \dots \times [365 - (N-1)]/365$. O primeiro termo deriva da opção de deixar a pessoa escolher livremente a data de aniversário (qualquer um dos 365 dias do ano é válido). O segundo termo corresponde à fração designada para a segunda pessoa se a sua data de aniversário não for igual à da primeira – pelo que tem a possibilidade de escolher entre 364 de 365 datas possíveis. O aniversário da terceira pessoa só pode calhar num dos 363 dos 365 dias restantes, para evitar ter a mesma data que uma das duas anteriores, etc, até chegarmos à n ésima pessoa que só pode escolher entre 365 menos $(N-1)$ dos 365 dias do ano, para evitar ter uma data de aniversário igual à de qualquer uma das $N-1$ pessoas anteriores. Desta forma, a probabilidade de duas delas fazerem realmente anos no mesmo dia é apenas $1 - P = 1 - N!/[365N(365 - N)!]$ que é superior a 0,5 quando N é maior do que 22.

10 Se alterarmos a probabilidade de confiança de 95% para outro número qualquer, por exemplo, $P\%$, os três intervalos do diagrama têm comprimentos de $\frac{1}{3} \times (1-P/100)$, $P/100$ e $\frac{1}{3} \times (1-P/100)$, respetivamente (verifique se o valor anterior dos intervalos é obtido com $P = 95$). Seguindo a mesma linha de raciocínio, vemos que no ponto A o futuro é $[P/100 + \frac{1}{3} \times (1-P/100)] / [\frac{1}{3} \times (1-P/100)] = [100 + P]/[100 - P]$ vezes mais longo do que o passado. Então, com $P\%$ de probabilidade, algo durará pelo menos $[100 - P]/[100 + P]$ vezes mais do que o seu tempo de vida atual, mas terá uma duração máxima de $[100 + P]/[100 - P]$ vezes esse tempo de duração. À medida que P se aproxima de 100%, as previsões começam a ser progressivamente menos exatas, por requererem um nível de exatidão muito maior. Se quisermos encontrar uma probabilidade de 99%, a duração prevista será 1/199 vezes a idade atual, mas menos de 199 vezes essa idade. Por outro lado, se deixarmos o nosso nível de certeza baixar para 50%, por exemplo, o intervalo provável é de 1/3 para 3 vezes a idade presente; é um intervalo muito estreito, mas a probabilidade de estar correto é demasiado pequena para ser preocupante.

11 Uma série Maclaurin é uma aproximação de série de uma função f de uma variável x por uma série de termos polinómiais. Conforme exigido, Tamm deverá ter escrito $f(x) = f(0) + x f'(0) + \dots + x^n f^{(n)}(0)/n! + R_n$ em que R_n é o erro, ou valor remanescente, depois de aproximar $f(x)$ por n dos termos anteriores, em que $n = 1, 2, 3, \dots$. O termo remanescente que foi pedido a Tamm que deduzisse pode ser calculado como sendo $R_n = \int_0^x x^{n+1} f^{(n+1)}(x)/n! dx$. Se o miliciano fosse muito precioso, podia ter exigido a simplificação adicional da expressão que é possível usando o teorema do valor médio para obter $R_n = x^{n+1} f^{(n+1)}(y)/(n+1)!$, em que y corresponde a um valor entre 0 e x . Colin Maclaurin era um Escocês contemporâneo de Isaac Newton.

12 Esta aproximação é mais adequada para valores muito baixos da taxa de juro j . Se quisermos torná-la mais exata, podemos fazer a aproximação de $\log(1+j)$ por $j - j^2/2$ e obter a estimativa de que $n = 0,7/j(1 - 1/2j)$. No caso das taxas de juro de 5%, em que $j = 0,05$, traduz o número de anos necessários para duplicar o investimento por 14,36 em vez de 14 anos.

13 Esta é uma belíssima aplicação daquilo a que em física se chama "método das dimensões". Taylor quer conhecer o raio, R , de uma explosão esférica num momento do tempo t após a detonação (chamemos ao tempo da detonação $t=0$). Supõe-se que depende da energia, E , liberada pela bomba e a densidade inicial, ρ , do ar circundante. Se existir uma fórmula $R = k E^a \rho^b t^c$, em que k, a, b e c são números a determinar, então, tendo em conta que as dimensões da energia são $M C^2 T^2$, a densidade $M C^3$, em que M é a massa, C é o comprimento e T é o tempo, o resultado terá de ser $a = 1/5$, $b = -1/5$ e $c = 2/5$. Assim, a fórmula é $R = k E^{1/5} \rho^{-1/5} t^{2/5}$. Supondo que a constante k é relativamente

próxima de 1, verificamos que o valor desconhecido da energia é determinado aproximadamente por $E = \mu R^2/r^2$. Comparando várias fotografias, também será possível determinar o valor de k .

¹⁴ É um elefante!

¹⁵ Qualquer número de 3 dígitos ABC pode ser escrito como $100A + 10B + C$. O primeiro passo consiste em subtrair o seu inverso, que é $100C + 10B + A$. O resultado é $99|A-C|$, em que as linhas verticais representam o valor positivo absoluto da diferença. Agora, a magnitude de $A-C$ deverá ser correspondente a um valor entre 2 e 9, pelo que ao multiplicá-lo por 99, obterá um dos múltiplos de 99 com três dígitos. Só há 8 respostas possíveis: 198, 297, 396, 495, 594, 693, 792 e 891. Repare que o dígito do meio é sempre 9 e que os outros dois somam sempre 9 em cada um dos números. Desta forma, ao somar qualquer um destes números ao seu inverso, obtém-se sempre o resultado 1089.

¹⁶ Se o líder disser a verdade com uma probabilidade p , a probabilidade de a sua afirmação ser verdadeira é $Q = p^2/[p^2 + (1-p)^2]$. No nosso exemplo, em que $p = 1/4$, obtemos uma probabilidade de $Q = 1/10$. Pode verificar que Q é menor do que p sempre que $p < 1/2$ e Q é maior do que p sempre que $p > 1/2$. Repare que quando $p = 1/2$, $Q = 1/2$.

¹⁷ Para ler uma explicação mais detalhada, consulte J. Haigh, *Taking Chances* (Correr Riscos), Oxford UP, Oxford (1999), capítulo 2. A quantidade essencial para calcular as probabilidades de ter r dos números premiados é 1 em ${}^{nC_r}/{}^{nC_r} = C_{n-r}/C_n$, em que $C_n = n!/(n-r)!r!$ é o número de diferentes combinações de r números distintos que podem ser seleccionadas a partir de n números.

¹⁸ Se criar uma sequência de Fibonacci (em que cada número a partir do terceiro é a soma dos dois anteriores), 1, 1, 2, 3, 5, 8, 13, 21, 34, 55, 89, etc. até ao infinito, designando o enésimo número da sequência como F_n , poderá generalizar este exemplo distribuindo um quadrado $F_n \times F_n$ no interior de um retângulo $F_{n+1} \times F_{n+1}$. O exemplo aqui representado é o caso $n = 6$ uma vez que $F_6 = 8$. A diferença das áreas do quadrado e do retângulo é determinada pela identidade de Cassini $(F_n \times F_n) - (F_{n+1} \times F_{n-1}) = (-1)^{n+1}$. Se n for um número par, o lado direito é igual a -1 e o retângulo é maior do que o quadrado; mas se n for um número ímpar, "perdemos" uma área de 1 quando mudamos da forma do quadrado para a do retângulo. Incrivelmente, a diferença é sempre +1 ou -1, independentemente do valor de n . Este enigma foi inventado por Lewis Carroll e tornou-se muito popular na Inglaterra vitoriana. Ver S. D. Collingwood, *The Lewis Carroll Picture Book*, Fisher Unwin, Londres (1899) pp. 316-7.

¹⁹ Para mais informações, ver Donald Saari, "Suppose You Want to Vote Strategically", *Maths Horizons*, novembro 2000, p. 9.

²⁰ Tolkowsky mostrou que para o reflexo interno total da luz que entra no diamante ocorrer na primeira face em que ela incide as facetas inclinadas têm de formar um ângulo de mais de $48.^\circ 52'$ em relação à horizontal. Depois do primeiro reflexo interno, a luz incidirá na segunda faceta inclinada e será completamente refletida se essa face formar um ângulo de menos de $43.^\circ 43'$ em relação à horizontal. Para a luz sair do diamante numa linha aproximadamente vertical (em vez da linha horizontal da face do diamante) e para ocorrer a melhor dispersão possível das cores da luz que está a sair, descobriu-se que o melhor valor para o ângulo era $40.^\circ 45'$. Os cortes modernos podem desviar-se ligeiramente destes valores para respeitarem as propriedades de cada pedra individual e para obter certas variações de estilo.

²¹ Para quaisquer duas quantidades x e y , é sempre verdade que $\frac{1}{2}(x+y) \geq (xy)^{1/2}$, que é uma consequência do facto de o quadrado $(x-y)^2 \geq 0$. Esta fórmula também apresenta uma característica muito apelativa que a média geométrica tem, mas que a média aritmética não. Se as quantidades x e y forem medidas em diferentes unidades (\$ por onça e £ por quilo, por exemplo) não fará grande sentido comparar o valor de $\frac{1}{2}(x+y)$ em momentos diferentes. Por o fazer, teria de se certificar de que x e y foram expressos nas mesmas unidades. No entanto, o valor da quantidade da média geométrica $(xy)^{1/2}$ pode ser comparado em diferentes momentos, mesmo que estejam a ser usadas unidades diferentes para medir x e y .